

l'intègre

Sous la direction de

Marie-Noëlle Sanz

Anne-Emmanuelle Badel

François Clausset

PHYSIQUE

TOUT-EN-UN • I^{re} année

MPSI - PCSI - PTSI

**NOUVELLE
ÉDITION**

- ▶ Un cours complet
- ▶ De nombreux exercices et problèmes
- ▶ Toutes les solutions détaillées

3^e édition

DUNOD

PHYSIQUE

TOUT-EN-UN • I^{re} année

Cours et exercices corrigés

MPSI - PCSI - PTSI

Sous la direction de

Marie-Noëlle Sanz

Professeur en PC au lycée Saint-Louis à Paris*

Anne-Emmanuelle Badel

*Professeur en BCPST
au lycée du Parc à Lyon*

François Clausset

*Professeur en MP
au lycée Jean Perrin à Lyon*

3^e édition

Couverture : *Bruno Loste*

Le pictogramme qui figure ci-contre mérite une explication. Son objet est d'alerter le lecteur sur la menace que représente pour l'avenir de l'écrit, particulièrement dans le domaine de l'édition technique et universitaire, le développement massif du photocopillage.

Le Code de la propriété intellectuelle du 1^{er} juillet 1992 interdit en effet expressément la photocopie à usage collectif sans autorisation des ayants droit. Or, cette pratique s'est généralisée dans les établissements

d'enseignement supérieur, provoquant une baisse brutale des achats de livres et de revues, au point que la possibilité même pour

les auteurs de créer des œuvres nouvelles et de les faire éditer correctement est aujourd'hui menacée. Nous rappelons donc que toute reproduction, partielle ou totale, de la présente publication est interdite sans autorisation de l'auteur, de son éditeur ou du

Centre français d'exploitation du droit de copie (CFC, 20, rue des Grands-Augustins, 75006 Paris).



© Dunod, Paris, 2002, 2003, 2008
ISBN 978-2-10-053977-2

Le Code de la propriété intellectuelle n'autorisant, aux termes de l'article L. 122-5, 2° et 3° a), d'une part, que les « copies ou reproductions strictement réservées à l'usage privé du copiste et non destinées à une utilisation collective » et, d'autre part, que les analyses et les courtes citations dans un but d'exemple et d'illustration, « toute représentation ou reproduction intégrale ou partielle faite sans le consentement de l'auteur ou de ses ayants droit ou ayants cause est illicite » (art. L. 122-4).

Cette représentation ou reproduction, par quelque procédé que ce soit, constituerait donc une contrefaçon sanctionnée par les articles L. 335-2 et suivants du Code de la propriété intellectuelle.

Table des matières



PREMIÈRE PÉRIODE

Électrocinétique 1

1

Chapitre 1 Lois générales de l'électrocinétique dans le cadre de l'approximation des régimes quasi-stationnaires

3

- | | | |
|---|---|----|
| 1 | Mouvement des porteurs de charges | 3 |
| 2 | Le courant électrique | 7 |
| 3 | Tension et potentiel | 10 |
| 4 | Loi des nœuds – loi des mailles | 11 |
| 5 | Notion de dipôles | 15 |
| 6 | Puissance – dipôles récepteurs et générateurs | 15 |

Chapitre 2 Circuits linéaires dans l'approximation des régimes quasi-stationnaires

18

- | | | |
|---|---|----|
| 1 | Dipôles linéaires | 18 |
| 2 | Résistor de résistance R | 19 |
| 3 | Bobine d'inductance L | 23 |
| 4 | Condensateur de capacité C | 25 |
| 5 | Sources de tension et de courant – Modèles de Thévenin et de Norton | 29 |
| 6 | Lois de Kirchhoff | 33 |
| 7 | Diviseurs de tension et de courant | 37 |
| 8 | Utilisation de l'équivalence entre les modèles de Thévenin et de Norton | 40 |
| A | Applications directes du cours | 43 |
| B | Exercices et problèmes | 46 |

Chapitre 3 Circuits linéaires soumis à un échelon de tension

53

- | | | |
|---|---------------------------|----|
| 1 | Définitions | 53 |
| 2 | Circuits du premier ordre | 55 |

Table des matières

3	Circuits du second ordre	67
A	Applications directes du cours	76
B	Exercices et problèmes	77

Mécanique 1 83

Chapitre 4	Cinématique du point matériel	85
1	Définitions	85
2	Systèmes de coordonnées planes	87
3	Description cinématique du mouvement d'un point matériel	92
4	Expressions de la vitesse et de l'accélération	94
5	Exemples de mouvements	96
A	Applications directes du cours	100
B	Exercices et problèmes	101
Chapitre 5	Principes de la dynamique newtonienne	104
1	Éléments cinétiques d'un point matériel	104
2	Notion de force	105
3	Les trois lois de Newton de la dynamique	112
4	Résolution d'un problème de mécanique - Exemples	116
A	Applications directes du cours	134
B	Exercices et problèmes	135
Chapitre 6	Aspects énergétiques de la dynamique du point	143
1	Travail et puissance d'une force	143
2	Théorème de l'énergie cinétique	145
3	Énergie potentielle et forces conservatives	148
4	Énergie mécanique	151
5	Application à l'étude qualitative des mouvements et des équilibres	152
6	Approche à l'aide des portraits de phase	155
7	Étude énergétique d'une masse suspendue à un ressort	162
A	Applications directes du cours	164
B	Exercices et problèmes	165
Chapitre 7	Oscillateurs harmoniques et amortis par frottement fluide	171
1	Oscillateurs harmoniques	171
2	Oscillateurs amortis	176

A	Applications directes du cours	187
B	Exercices et problèmes	188
Optique		195
Chapitre 8	Approximation de l'optique géométrique, rayon lumineux	197
1	Notion expérimentale de rayon lumineux	197
2	Aspect ondulatoire	198
3	Lois fondamentales de l'optique géométrique	200
Chapitre 9	Réfraction, réflexion	203
1	Le vocabulaire de l'optique géométrique	203
2	Lois de Descartes	206
3	Cas des milieux stratifiés. Mirages	211
A	Applications directes du cours	214
B	Exercices et problèmes	215
Chapitre 10	Stigmatisme et aplanétisme. Dioptries et miroirs	221
1	Stigmatisme et aplanétisme rigoureux	221
2	Stigmatisme et aplanétisme approchés	223
3	Cas du miroir sphérique	227
A	Applications directes du cours	241
B	Exercices et problèmes	242
Chapitre 11	Lentilles minces sphériques	247
1	Présentation	247
2	Relations de conjugaison	250
3	Construction	253
4	Quelques instruments d'optique	260
5	Projection d'image et aberrations	264
A	Applications directes du cours	271
B	Exercices et problèmes	273
Chapitre 12	Instruments de Travaux Pratiques	278
1	L'œil	278
2	Sources de lumière	282
3	Loupes et oculaires	285
4	Lunettes et viseurs	287

Table des matières

5	Collimateur	293
6	Le goniomètre	293
A	Applications directes du cours	301
B	Problèmes	301
Chapitre 13 T.P. : Focométrie		304
1	Reconnaissance rapide du caractère d'une lentille ou d'un miroir	304
2	Méthode directe	306
3	Autocollimation	308
4	Méthode de Silberman	309
5	Méthode de Bessel	310
6	Méthode des points conjugués	311
Chapitre 14 T.P. : Spectroscopie à prisme		313
1	Présentation	313
2	Application des lois de Descartes au prisme	313
3	Spectromètre à prisme	318

DEUXIÈME PÉRIODE

Électrocinétique 2 321

Chapitre 15 Circuits linéaires en régime sinusoïdal forcé		323
1	Caractéristiques d'un signal sinusoïdal	323
2	Intérêt de la notation complexe	326
3	Notation complexe	328
4	Lois de Kirchhoff en notation complexe	334
5	Impédance et admittance complexes	336
6	Loi des nœuds exprimée en termes de potentiels	342
7	Régime sinusoïdal forcé et régime transitoire	344
8	Stabilité des circuits du premier et du second ordre (PCSI)	347
A	Applications directes du cours	350
B	Exercices et problèmes	351
Chapitre 16 Circuit R,L,C série en régime sinusoïdal forcé et résonances		354
1	Circuit R,L,C série soumis à une excitation sinusoïdale	354
2	Résonance en intensité	357
3	Résonance aux bornes de la capacité	364

4	Étude de l'impédance	369
A	Applications directes du cours	372
B	Exercices et problèmes	372
Chapitre 17 Puissance		378
1	Définitions	378
2	Puissance en régime sinusoïdal forcé	382
A	Applications directes du cours	391
B	Exercices et problèmes	392
Chapitre 18 Réponse fréquentielle d'un circuit linéaire - Filtres linéaires du premier et du second ordre		395
1	Notion de quadripôle	395
2	Fonction de transfert	396
3	Gain	399
4	Étude du comportement en fonction de la fréquence. Diagramme de Bode	401
5	Filtres du premier ordre	404
6	Filtres du deuxième ordre (PCSI, PTSI)	417
7	Lien entre fonction de transfert et équation différentielle. Transformations réciproques	436
A	Applications directes du cours	439
B	Exercices et problèmes	442
Chapitre 19 Étude expérimentale de quelques filtres		447
1	Décomposition harmonique d'un signal	447
2	Analyse harmonique et circuits linéaires	455
3	Notion de filtrage	457
4	Résultats du filtrage	460
5	Étude d'un filtre sélectif	468
A	Applications directes du cours	472
B	Exercices et problèmes	472
Chapitre 20 T.P. Cours : Instrumentation électrique		478
1	Générateurs basses fréquences ou GBF (PCSI, PTSI)	478
2	Alimentations stabilisées (PCSI, PTSI)	481
3	Oscilloscopes	482
4	Masse et Terre	490

Table des matières

5	Multimètres	493
6	Modèles réels des dipôles linéaires passifs	495
7	Mesure de déphasages	500
8	Mesures d'impédances par diviseur de tension	502
A	Applications directes du cours	505
B	Exercices et problèmes	505
Chapitre 21	T.P. Cours : Amplificateur opérationnel	508
1	Modèle de l'amplificateur idéal	508
2	Deux montages simples	512
3	Quelques écarts au modèle idéal de l'amplificateur opérationnel	514
4	Montage suiveur	523
5	Montage amplificateur non inverseur	528
6	Montages intégrateur et pseudo-intégrateur	532
7	Complément : autres exemples de montages simples	536
8	Comparateurs à amplificateur opérationnel	539
9	Complément : multivibrateur astable	544
A	Applications directes du cours	549
B	Exercices et problèmes	552
Chapitre 22	Étude expérimentale de quelques circuits à diode (PCSI)	563
1	La diode : un dipôle non linéaire	563
2	Redressement et filtrage	569
3	Enrichissement du spectre par un dipôle non linéaire	577
4	Multiplieurs analogiques	578
A	Applications directes du cours	581
B	Exercices et problèmes	583

Mécanique 2

593

Chapitre 23	Compléments de cinématique du point matériel	595
1	Systèmes de coordonnées dans l'espace	595
2	Dérivée d'un vecteur unitaire tournant par rapport à l'angle de rotation	604
3	Expressions de la vitesse et de l'accélération	605
4	Compléments sur le mouvement circulaire	607
A	Applications directes du cours	609

B	Exercices et problèmes	609
Chapitre 24	Oscillations mécaniques forcées	610
1	Position du problème et équation du mouvement	610
2	Étude de la solution	611
3	Étude de l'amplitude en fonction de la fréquence - Résonance en amplitude	614
4	Résonance en vitesse	617
5	Impédance mécanique	619
6	Aspects énergétiques - Résonance en puissance	620
7	Portrait de phase d'un oscillateur forcé	622
A	Applications directes du cours	624
B	Exercices et problèmes	625
Chapitre 25	Théorème du moment cinétique	631
1	Définition du moment cinétique d'un point matériel	631
2	Moment d'une force	633
3	Théorème du moment cinétique en référentiel galiléen	634
4	Exemple d'utilisation du théorème du moment cinétique : le pendule simple	637
A	Applications directes du cours	640
B	Exercices et problème	641
Chapitre 26	Changement de référentiel : aspects cinématiques	644
1	Position du problème	644
2	Composition des vitesses	647
3	Composition des accélérations	650
A	Applications directes du cours	654
B	Exercices et problèmes	655
Chapitre 27	Changement de référentiel : aspects dynamiques	659
1	Relativité galiléenne	659
2	Expression du principe fondamental de la dynamique en référentiel non galiléen	663
3	Théorème du moment cinétique en référentiel non galiléen	665
4	Théorème de l'énergie cinétique en référentiel non galiléen	666
5	Exemple de résolution d'un problème dans un référentiel non galiléen	667

Table des matières

A	Application directe du cours	669
B	Exercices et problèmes	671
Chapitre 28	Cinétique d'un système de deux points matériels	677
1	Masse et centre d'inertie	677
2	Éléments dynamiques d'un système de points matériels	678
3	Référentiel barycentrique et théorèmes de Koenig	680
Chapitre 29	Dynamique d'un système de deux points matériels	684
1	Actions extérieures et intérieures	684
2	Théorème de la résultante cinétique	686
3	Théorème du moment cinétique	688
4	Puissance et travail d'un système de forces	690
5	Théorème de l'énergie cinétique - Énergie mécanique	693
A	Application directe du cours	699
B	Exercices et problèmes	699
Chapitre 30	Systèmes de deux points matériels isolés	702
1	Décomposition du mouvement	702
2	Référentiel barycentrique et éléments dynamiques	703
3	Particule fictive	706
4	Lois de conservation et conséquences	707
5	Analyse qualitative du mouvement	711
A	Application directe du cours	715
B	Exercices et problèmes	715
Chapitre 31	Mouvement à force centrale et potentiel newtonien	721
1	Force centrale conservative	721
2	Potentiel newtonien	726
3	Détermination de l'expression de la trajectoire (MPSI, PCSI)	728
4	Discussion sur la nature du mouvement (MPSI, PCSI)	734
5	Mouvement des planètes - Lois de Kepler	736
6	Quelques remarques sur le mouvement des satellites (MPSI, PCSI)	743
A	Applications directes du cours	748
B	Exercices et problèmes	752
Chapitre 32	Quelques aspects de la mécanique terrestre (PCSI)	759
1	Interaction gravitationnelle	759

2	Champ de pesanteur terrestre	760
3	Terme de marée	766
4	Exemple de l'influence de la force de Coriolis : déviation vers l'est	773
A	Application directe du cours	777
B	Exercices et problèmes	778
Thermodynamique		785
Chapitre 33 Généralités sur les systèmes thermodynamiques		786
1	Notions de base	786
2	Température	790
3	Équation d'état	795
A	Applications directes du cours	800
B	Exercices et problèmes	801
Chapitre 34 Statique des fluides dans le champ de pesanteur		803
1	Définition	803
2	Pression d'un fluide soumis au champ de pesanteur	804
3	Actions exercées par les fluides au repos. Poussée d'Archimède	811
A	Applications directes du cours	816
B	Exercices et problèmes	817
Chapitre 35 Interprétation microscopique de T et P. Énergie interne. Cas du gaz parfait monoatomique		822
1	Les hypothèses	822
2	Calcul de la pression	825
3	Température cinétique	827
4	Énergie interne	828
A	Applications directes du cours	833
B	Exercices et problèmes	834
Chapitre 36 Premier principe		835
1	Premier principe – Énergie interne	835
2	Transformations d'un système thermodynamique	839
3	Expression du travail	842
4	Enthalpie	846

Table des matières

Chapitre 37 Applications du premier principe	851
1 Calorimétrie	851
2 Transformations d'un gaz parfait	855
3 Travail fourni par un opérateur lors du déplacement d'un piston	860
4 Détente de Joule-Gay Lussac	861
5 Détente de Joule - Thomson ou Joule - Kelvin	864
6 Récapitulatif : utiliser U ou H	868
A Applications directes du cours	869
B Exercices et problèmes	871
Chapitre 38 Deuxième principe de la thermodynamique	877
1 Nécessité d'un second principe. Limites du premier principe	877
2 Réversibilité - Irréversibilité	878
3 Température et pression thermodynamique	881
4 Exemples de calculs d'entropie	885
5 Relation de Clausius	896
A Applications directes du cours	898
B Exercices et problèmes	901
Chapitre 39 Interprétation statistique de l'entropie (PCSI)	905
1 Microétat et macroétat	905
2 État macroscopique le plus probable	906
3 Entropie statistique	909
4 Troisième principe de la thermodynamique	911
5 Complément : entropie et information	912
Chapitre 40 Équilibre d'un corps pur sous plusieurs phases	913
1 Généralités	913
2 Diagramme d'équilibre	916
3 Les fonctions d'état	925
4 Expression des fonctions d'états à partir de valeurs tabulées	931
A Applications directes du cours	932
B Exercices et problèmes	934
Chapitre 41 Machines thermiques	938
1 Description	938
2 Machine monotherme	939

3	Étude d'une machine ditherme	940
4	Exemples de machines dithermes	943
5	Autres exemples de machines thermiques	949
A	Applications directes du cours	957
B	Exercices et problèmes	958

Électromagnétisme 965

Chapitre 42	Charge électrique et distributions de charges	966
1	Distributions de charges	966
2	Déplacements, surfaces et volumes élémentaires	970
3	Invariances par transformation géométrique	972
Chapitre 43	Champ électrostatique	976
1	Loi de Coulomb	976
2	Notion de champ	978
3	Définition du champ électrostatique	979
4	Analogie avec le champ de gravitation	981
5	Invariances, symétries et propriétés du champ	984
6	Application au champ électrostatique créé par un segment dans son plan médiateur ou par un fil infini uniformément chargé	990
A	Applications directes du cours	993
B	Exercices et problèmes	993
Chapitre 44	Propriétés du champ électrostatique	996
1	Flux du champ électrostatique - Théorème de Gauss	996
2	Circulation du champ électrostatique - Potentiel électrostatique	1003
3	Énergie électrostatique	1010
4	Topographies du champ et du potentiel électrostatiques	1016
A	Applications directes du cours	1020
B	Exercices et problèmes	1021
Chapitre 45	Exemples de champs et potentiels électrostatiques	1025
1	Méthodes de calcul	1025
2	Champ et potentiel créés par une sphère uniformément chargée en volume	1028
3	Champ et potentiel créés par un cylindre ou un fil uniformément chargé	1032
4	Champ et potentiel créés par un plan uniformément chargé	1036

Table des matières

5	Champ et potentiel créés par un disque uniformément chargé le long de son axe	1042
A	Applications directes du cours	1045
B	Exercices et problèmes	1047
Chapitre 46 Dipôle électrostatique (MPSI, PCSI)		1053
1	Définition, potentiel et champ créés	1053
2	Action d'un champ extérieur sur un dipôle	1058
3	Complément : approximation dipolaire	1066
A	Applications directes du cours	1072
B	Exercices et problèmes	1074
Chapitre 47 Courant électrique et distributions de courants		1078
1	Courant électrique	1078
2	Densité de courant	1080
3	Invariances d'une distribution de courants	1084
Chapitre 48 Champ magnétique créé par des courants permanents		1086
1	Définition du champ magnétique	1086
2	Propriétés de symétrie du champ magnétique et conséquences	1088
3	Loi de Biot et Savart	1092
4	Interactions magnétiques	1095
A	Applications directes du cours	1097
B	Exercices et problèmes	1098
Chapitre 49 Propriétés du champ magnétique		1099
1	Conservation du flux	1099
2	Circulation du champ magnétique - Théorème d'Ampère	1101
3	Topographie	1102
A	Applications directes du cours	1107
B	Exercices et problèmes	1108
Chapitre 50 Exemples de champs magnétiques		1110
1	Méthodes de calculs	1110
2	Fil rectiligne infini	1111
3	Spire circulaire	1114
4	Bobine torique	1120

Table des matières

5	Solénoïde	1121
6	Complément : nappe plane de courant	1128
A	Applications directes du cours	1131
B	Exercices et problèmes	1132
Chapitre 51	Dipôle magnétique (PCSI)	1138
1	Définitions	1138
2	Champ magnétique créé par un dipôle magnétique	1140
3	Complément : action d'un champ magnétique extérieur sur un dipôle magnétique	1142
A	Applications directes du cours	1145
B	Exercices et problèmes	1145
Chapitre 52	Comparaison des propriétés des champs électrostatique et magnétique	1147
1	Définitions des champs	1147
2	Flux	1147
3	Circulation	1147
4	Expression à partir de ses sources	1148
5	Propriétés de symétrie	1148
6	Récapitulatif des propriétés des champs électrostatique et magnétique	1150
7	Dipôles électrostatique et magnétique	1151
Chapitre 53	Mouvement d'une particule chargée dans un champ électrique ou magnétique	1152
1	Force de Lorentz	1152
2	Mouvement dans un champ électrique	1154
3	Mouvement d'une particule chargée dans un métal - modèle de la loi d'Ohm locale	1161
4	Mouvement dans un champ magnétique	1165
5	Mouvement dans un champ électrique et magnétique	1174
6	Effet Hall	1179
A	Applications directes du cours	1184
B	Exercices et problèmes	1186
Solutions des exercices		1192
Index		1516

Exercices corrigés supplémentaires sur www.dunod.com

Vous pouvez accéder à des exercices supplémentaires et leurs corrigés complets sur le site Internet www.dunod.com.

- Pour cela, entrez « Sanz » (le nom du premier auteur) dans le moteur de recherche.
- Une fois la recherche effectuée, cliquez sur l'ouvrage « Physique tout-en-un 1^{re} année MPSI-PCSI-PTSI » pour atteindre la fiche de présentation du livre.
- Cliquez ensuite sur « Compléments en ligne » pour accéder aux exercices corrigés supplémentaires.

PREMIÈRE PÉRIODE

Électrocinétique 1

Lois générales de l'électrocinétique dans le cadre de l'approximation des régimes quasi-stationnaires

1

L'électrocinétique concerne l'étude du mouvement de particules chargées dans la matière sous l'action d'un champ électrique. Dans ce chapitre, on définit les notions fondamentales comme le courant et la tension et on précise les lois générales de l'électricité dans le cadre de l'approximation des régimes quasi-stationnaires.

1. Mouvement des porteurs de charges

1.1 Notion de charge électrique

Certains corps sont susceptibles d'accepter ou de perdre des particules chargées : on dit qu'ils s'électrisent. On peut citer :

- le verre qui, frotté avec de la soie, perd des électrons,
- l'ébonite ou l'ambre qui, frottés avec une fourrure (par exemple une peau de chat), acquièrent des électrons.

L'électrisation obéit à plusieurs lois qualitatives :

- les corps électrisés exercent des actions mécaniques : ils attirent par exemple des objets légers comme des petits bouts de papier (c'est ce type d'expériences mettant en évidence l'électrisation des corps qui a permis la découverte de l'électricité dès l'Antiquité),
- l'électrisation peut se transférer d'un corps à un autre,
- il existe deux types d'électrisation qui seront qualifiés conventionnellement de positive et de négative,
- deux corps de même type d'électrisation se repoussent tandis que deux corps de type différent d'électrisation s'attirent,
- tout corps électrisé peut attirer un corps non électrisé.

Ces résultats expérimentaux s'interprètent par la notion de *charge électrique* obtenue grâce aux travaux de Coulomb de 1785 et à la découverte de l'électron en 1881 par Thomson. La charge électrique, qui caractérise le phénomène d'électrisation de la charge élémentaire, est indissolublement liée à la matière.

La charge électrique est une grandeur scalaire positive ou négative vérifiant les propriétés suivantes.

a) Charge positive ou négative

Elle peut exister sous deux formes qu'on qualifie de positive et de négative. Le choix de l'électron comme porteur d'une charge négative est purement conventionnel mais admis de tous. Une charge sera donc positive si elle est attirée par un électron et négative si elle est repoussée par ce dernier. Ceci permet de satisfaire les phénomènes d'attraction et de répulsion observés expérimentalement.

b) Extensivité de la charge

La charge électrique est une grandeur extensive dans la même acception que celle qui sera adoptée en thermodynamique : la charge ne dépend que de l'état du système et elle est égale à la somme algébrique des charges élémentaires qui la constituent.

On peut formuler cette définition en considérant un système formé de l'association de deux sous-systèmes, l'un de charge électrique q_1 , l'autre de charge électrique q_2 . La charge q du système est :

$$q = q_1 + q_2$$

c) Conservation de la charge

La charge électrique est une grandeur conservative au sens où la charge électrique totale d'un système isolé est constante au cours du temps.

Les variations de la charge d'un système ne sont donc dues qu'aux échanges avec l'extérieur. Pour un système isolé, il n'y a ni création ni disparition de charges, sa charge électrique ne varie pas.

d) Existence d'une charge élémentaire – Quantification de la charge

Les lois de l'électrolyse découvertes par Michael Faraday (1791-1862) s'interprètent par l'existence d'une charge électrique élémentaire notée e qui vaut $1,6 \cdot 10^{-19}$ C ou coulombs. Le coulomb est l'unité de charge.

L'ensemble des expériences qui ont été réalisées à ce jour indique que toute charge électrique rencontrée dans la nature est un multiple entier de cette charge, ce qui justifie le fait de parler de charge élémentaire :

$$q = \pm Ze \quad \text{avec} \quad Z \in \mathbb{N}$$

On introduit ainsi la notion de quantification de la charge électrique : celle-ci ne peut prendre qu'un certain nombre de valeurs qui sont les multiples de la charge élémentaire. Cette idée de la quantification de la charge électrique est apparue lors de la découverte de la structure de l'atome. La charge élémentaire joue donc le rôle de quantum pour les charges électriques.

► **Remarque** : À l'échelle macroscopique¹, la charge apparaît continue comme on le verra dans le cours d'électromagnétisme.

e) Invariance de la charge

La charge électrique est invariante par changement de référentiel galiléen² : quel que soit le référentiel dans lequel on se place, la charge a toujours les mêmes caractéristiques.

f) Convention de la charge

Le fait d'attribuer un signe + ou un signe - à une charge ou à un porteur de charge est une pure convention mais il est important de s'y conformer afin que tout le monde parle le même langage.

1.2 Porteurs de charges

a) Rappel sur les atomes

La matière est constituée d'atomes qui sont formés :

- d'un noyau comportant :
 - des neutrons, particules non chargées, de masse $m_n = 1,675 \cdot 10^{-27}$ kg,
 - des protons, particules chargées positivement par convention, leur charge étant la valeur de la charge élémentaire $q_p = e = 1,6 \cdot 10^{-19}$ C et de masse $m_p = 1,673 \cdot 10^{-27}$ kg,
- d'un cortège électronique constitué d'électrons qui sont des particules chargées négativement par convention et dont la charge est en valeur absolue la charge élémentaire $q_e = -e = -1,6 \cdot 10^{-19}$ C et de masse $m_e = 9,109 \cdot 10^{-31}$ kg

Les atomes sont électriquement neutres : leur charge totale est nulle et il y a donc autant de protons que d'électrons dans les atomes.

b) Conduction métallique

Les métaux sont des empilements réguliers d'atomes. La conduction métallique est liée à l'existence d'électrons dits libres ou électrons de conduction. Le phénomène

¹ Cette notion d'échelle sera reprise dans le cours de thermodynamique et d'électromagnétisme.

² Voir cours de mécanique.

s'explique dans le cadre de la théorie des bandes qui est hors programme. Il suffit ici de savoir qu'en moyenne un atome du métal conducteur comme le cuivre libère un électron de conduction. Ce dernier peut se déplacer³ au sein du métal entre les différents atomes : il n'est pas lié à un atome d'où le nom de libre qui leur est donné. Les porteurs de charge dans un conducteur métallique sont les électrons libres.

c) Solutions électrolytiques

Les porteurs de charge sont dans ce cas les ions. Les atomes décrits précédemment peuvent céder ou gagner des électrons : on parle d'ionisation de l'atome. Les ions ainsi obtenus ont donc par définition une charge non nulle contrairement aux atomes.

On distingue deux types d'ions en fonction du signe de leur charge totale :

- les cations qui ont une charge positive, ce qui correspond à une perte d'électrons,
- les anions qui ont une charge négative, ce qui correspond à un gain d'électrons.

La charge des ions est un multiple de la charge élémentaire suivant le nombre d'électrons qui ont été gagnés ou cédés.

Il convient de noter que la matière est globalement neutre conformément au principe de conservation de la charge. Le phénomène d'ionisation correspond à une modification de la répartition des charges mais il n'y a ni apparition ni disparition de charges.

d) Plasmas

Les plasmas sont des gaz ionisés à haute température composés d'ions positifs et d'électrons. Ces deux types de charges peuvent être les porteurs de charge. Cependant, le plus souvent, les ions sont considérés comme immobiles du fait de leur masse très supérieure à celle des électrons.

e) Semi-conducteurs

Ce type de matériau est très particulier au niveau de la conduction. On ne s'intéresse ici qu'à la nature des porteurs de charges.

On distingue deux types de porteurs majoritaires :

- les électrons dans le cas des semi-conducteurs dopés N,
- les « trous » qui correspondent à des lacunes d'électrons ou à un manque local d'électrons autour de certains atomes dans le cas des semi-conducteurs dopés P. Ces « trous » se comportent comme des charges pouvant se déplacer comme les électrons : on aura donc un déplacement de charges positives.

³Nous reverrons ce phénomène au chapitre sur le mouvement des particules chargées lors de l'étude de la loi d'Ohm.

Globalement on a vu comme porteurs de charge :

- des électrons,
- des ions,
- des trous.

Nous allons maintenant nous intéresser à leur mouvement.

1.3 Mouvement d'agitation thermique ou mouvement d'ensemble – courant électrique

On verra dans le cours de thermodynamique que toute particule est soumise au phénomène d'agitation thermique. Il s'agit d'un mouvement désordonné aléatoire qui reste local : il n'y a pas de mouvement d'ensemble au sens où les particules se déplaceraient toutes de la même manière. Les porteurs de charges comme toute particule microscopique sont soumis à ce type de mouvement.

À ce mouvement d'agitation thermique peut se superposer un mouvement d'ensemble sous une action extérieure dont l'origine peut être diverse : champ de gravitation, champ électrique, etc. Sous cette action extérieure, toutes les particules subissent la même force et se déplacent par conséquent de la même manière. Lorsque les particules sont des porteurs de charges, leur mouvement d'ensemble provoque un déplacement de charges électriques qu'on appelle *courant électrique*.

En électrocinétique, on ne considère cependant pas tous les courants électriques : on ne s'intéresse qu'à ceux dont l'origine du mouvement d'ensemble est de nature électrique. On se limite donc par la suite au cas où les porteurs de charges sont mis en mouvement sous l'action d'un champ électrique.

2. Le courant électrique

2.1 Définition

Soit un fil de section S quelconque. On soumet ce fil à l'action d'un champ électrique extérieur orienté le long de ce fil. On admet arbitrairement que l'orientation du champ électrique oriente le fil :

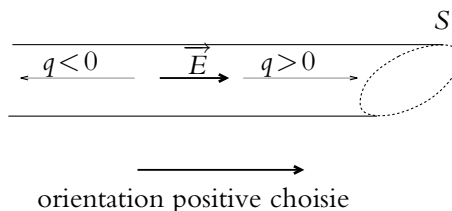


Figure 1.1 Définition du courant électrique.

Sous l'action du champ électrique extérieur, les porteurs de charges sont soumis à la force $\vec{f} = q\vec{E}$ et sont donc animés d'un mouvement d'ensemble tel que :

- les charges positives se déplacent dans le sens du champ,
- les charges négatives dans le sens contraire.

Au sens défini précédemment, il existe un courant électrique auquel on s'intéresse en électrocinétique.

On peut noter qu'un courant de charges négatives se déplaçant dans un sens est équivalent à un courant de charges opposées se déplaçant dans l'autre sens. Ainsi, dans une électrolyse, puisque les cations et les anions ont des charges de signe opposé et qu'ils se déplacent en sens opposé, les courants s'ajoutent au lieu de se compenser.

2.2 Sens conventionnel du courant

On a choisi le sens conventionnel du courant avant de découvrir qu'il est assuré par les électrons libres dans les conducteurs métalliques, ce qui explique que le sens conventionnel ne correspond pas à celui du déplacement des électrons. En effet, par convention, **le sens du courant est le sens dans lequel se déplaceraient les charges positives soumises au champ électrique extérieur**. Une autre façon de le définir consiste à dire que le sens conventionnel du courant est le sens inverse du mouvement des électrons sous l'action d'un champ électrique extérieur.

L'avantage d'une telle convention est de ne pas avoir à connaître la nature des charges et notamment leur signe pour étudier le courant électrique qui en résulte. Il s'agit d'un choix tout aussi conventionnel que celui de l'orientation des axes pour se repérer dans l'espace.

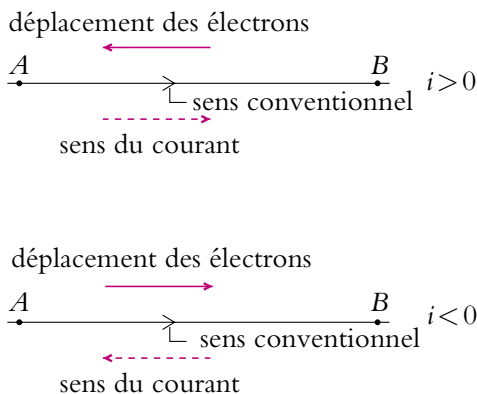


Figure 1.2 Sens conventionnel, sens du courant et déplacement des électrons.

2.3 Signe du courant

En mécanique, avant d'étudier un mouvement et de savoir si un mobile va se déplacer vers la droite ou vers la gauche, on définit l'orientation des axes. C'est la même chose

en électricité pour le sens du courant. Avant d'effectuer tout calcul sur un circuit, on ne sait pas *a priori* quel sera le sens du courant dans un fil donné. Il faut donc **choisir un sens arbitraire** pour le courant. Par exemple, dans le cas précédent, si l'on ne sait pas que le courant se déplace de B vers A , on peut choisir un sens positif du courant de A vers B . On représente cette orientation avec une flèche (et éventuellement un signe $+$ à côté).

Avec ce choix, si effectivement le courant circule de A vers B , il est positif comme sur la figure du haut, mais s'il circule de B vers A , il est négatif.

► **Remarque :** En réfléchissant avant un calcul, on peut parfois avoir une idée du sens du courant dans un circuit simple, on préférera alors choisir le sens du courant dans le sens attendu.

2.4 Intensité du courant

Soit un conducteur comme celui représenté sur la figure 2.1. On note S la section de ce conducteur.

Entre t et $t + \delta t$, la quantité de charges δq traverse la section S du conducteur du fait de l'existence du courant électrique. La charge δq est une grandeur algébrique par rapport au sens positif choisi. Si des charges positives se déplacent effectivement dans ce sens, δq est positive ; si elles se déplacent effectivement dans le sens opposé au sens positif conventionnellement choisi, δq est négative. Inversement, si des charges négatives se déplacent effectivement dans le sens positif conventionnellement choisi, δq est négative ; si elles se déplacent effectivement dans le sens opposé au sens positif conventionnellement choisi, δq est positive.

On appelle *intensité du courant électrique* et on note i la quantité de charges traversant S par unité de temps. Cela se traduit mathématiquement par : $\delta q = i \delta t$ soit :

$$i = \frac{\delta q}{\delta t}$$

L'unité de l'intensité est l'ampère et son symbole est A. L'ampère est une des unités de base du Système International. Cela correspond à des $C.s^{-1}$. Sur le schéma d'un circuit, on écrira i à côté de l'orientation conventionnelle choisie pour le fil considéré.

On notera que l'intensité i tout comme la charge électrique est une **grandeur algébrique** pouvant être positive ou négative. Si elle est constante au cours du temps, on dira que le courant est *continu*.

2.5 Quelques ordres de grandeur d'intensité

Pour avoir une idée des ordres de grandeur d'intensités utilisées par les appareils domestiques, le lecteur pourra consulter les étiquettes de fusibles d'une installation électrique.

Ainsi les fusibles sont de 16 A pour les prises électriques de courant, de 32 A pour un four ou des plaques électriques. Un chauffage de 1000 W consomme environ 5 A.

En électronique, les intensités d'entrée d'un circuit comme un amplificateur opérationnel sont inférieures (pour les plus récents) à 10^{-12} A, alors que le courant à la sortie peut avoir une intensité allant jusqu'à 20 mA. Les alimentations continues d'appareils électroniques peuvent délivrer des intensités de l'ordre de l'ampère.

Un T.G.V. consomme un courant de plusieurs centaines d'ampères (500 A en régime de croisière avec des pointes à 1000 A).

L'intensité dans les lignes de distribution électrique hautes tensions est de l'ordre de 1000 A.

Dans l'industrie, les ordres de grandeurs sont encore plus élevés. Dans les fours d'acier, l'intensité utilisée est de l'ordre de 10^5 A.

Les ordres de grandeur des intensités sont donc très variés.

3. Tension et potentiel

3.1 Définitions

On appelle *tension* ou *différence de potentiel* la grandeur mesurée par un voltmètre entre deux points *A* et *B*. Elle s'exprime en volt, de symbole V, en hommage au physicien Volta (1745 – 1827).

On notera les tensions avec la lettre *U* et les potentiels avec la lettre *V*. Ainsi la tension U_{AB} entre deux points A et B d'un conducteur est égale à la différence de potentiel entre ces deux points A et B :

$$U_{AB} = V_A - V_B$$

Aux bornes d'un élément de circuit, qu'on représente par un rectangle sur la figure 1.3, on mesure une tension *U*. On indique cette tension sur le schéma par une flèche dont le sens est très important. En effet, il s'agit du choix de l'orientation de la tension soit $U_1 = V_B - V_A$ soit $U_2 = V_A - V_B$. Ce choix, comme celui de l'orientation de l'intensité, est parfaitement arbitraire mais permet de déterminer le point dont le potentiel est le plus élevé. Ainsi si $U_1 > 0$ alors le point *A* a un potentiel plus élevé que le point *B*.

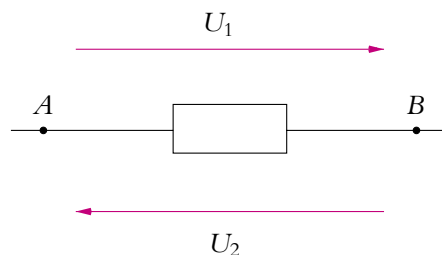


Figure 1.3 Représentation de la tension.

3.2 Masse ou référence de potentiel

La tension qu'on peut mesurer expérimentalement est une différence de potentiel entre deux points. Aucun appareil ne permet d'accéder à la mesure du potentiel en un point donné. Cela traduit expérimentalement le résultat, qu'on établira dans le cours d'électrostatique, que le potentiel en un point est défini à une constante près.

Pour fixer cette constante, on choisit arbitrairement une référence de potentiel nul, qu'on appelle la *masse*. Ainsi si on choisit par exemple comme masse une borne de l'oscilloscope, ce dernier donne la tension entre ses deux bornes.

Pour des raisons de sécurité, on relie la carcasse des appareils à la Terre. Souvent la Terre est également reliée à une borne de l'appareil : la masse est alors prise à la Terre. Les appareils pour lesquels cette liaison n'existe pas sont dits à *masse flottante*. On reviendra plus précisément sur ces deux notions dans le chapitre sur l'instrumentation électrique.



Figure 1.4 Symboles de la masse et de la Terre.

► **Remarque :** On retiendra qu'on parle de la tension **aux bornes** d'un élément d'un circuit et de l'intensité **traversant** un élément de circuit.

4. Loi des nœuds - loi des mailles

4.1 Terminologie des circuits

Avant d'étudier les circuits électriques, on a besoin de définir quelques termes relatifs à leur constitution.

- Un *dipôle* est un élément de circuit relié au reste du circuit par deux bornes.
- Une *branche* est un ensemble de dipôles reliés par des fils de connexion et disposés en série c'est-à-dire que chaque borne d'un dipôle n'est reliée qu'à un seul autre dipôle.

L'ensemble des éléments d'un circuit électrique est appelé un *réseau*.

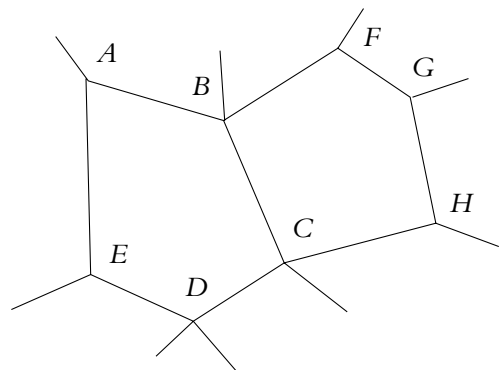


Figure 1.5 Partie d'un circuit électrique.

- Un *nœud* est un point où se rejoignent au moins deux branches.
- Une *maille* est un ensemble de branches se refermant sur elles-mêmes.

Sur la figure 1.5, on a représenté une portion de circuit. Les fils dont une extrémité est libre sur le schéma sont en fait reliés à une partie du réseau non représentée.

- $AB, BC, CD, DE, EA, BF, FG, GH$ et HC sont des branches.
- A, B, C, D, E, F, G et H sont des nœuds.
- $ABCDEA, BFGHCB$ et $ABFGHCDEA$ sont des mailles.

4.2 Régime continu ou variable

On dit qu'on est en *régime continu* lorsque toutes les grandeurs sont indépendantes du temps ; ce sera notamment le cas des intensités et des tensions.

A contrario, on parle de *régime variable* quand les grandeurs dépendent du temps. Le caractère variable peut avoir plusieurs origines possibles pouvant se combiner :

- modification des conditions extérieures faisant passer d'un régime continu à un autre : on parlera alors de régime transitoire,
- conditions extérieures variables par exemple de type sinusoïdales ou créneau : on parlera alors de régime forcé,
- phénomène de propagation : comme la pression lors de la propagation d'une onde sonore ou la lumière (Cf. cours de terminale), les tensions et intensités se propagent dans les conducteurs. Cela signifie que leur valeur dépend à la fois du temps et du point considéré. La vitesse de propagation de l'intensité et de la tension est celle de la lumière dans le vide à savoir $c = 3.10^8 \text{ m.s}^{-1}$. La durée de propagation dans un fil conducteur ou dans un circuit de longueur L peut s'estimer par : $\tau = \frac{L}{c}$. La durée τ est le temps caractéristique du phénomène de propagation de l'intensité et de la tension dans le circuit considéré.

4.3 Approximation des régimes quasi-stationnaires ou ARQS

Dans un circuit en régime continu, il n'y a pas d'accumulation de charges : l'intensité est donc la même en tout point d'une branche.

La mesure de l'intensité dans une branche avec un ampèremètre ne dépend alors pas de la position de l'appareil le long de la branche.

Cette propriété reste valable en régime variable si on peut négliger les phénomènes de propagation. Cela revient à considérer que le temps de propagation est très petit devant le temps caractéristique du régime variable. La propagation sera assimilable à un processus instantané et l'intensité dans une branche sera la même en tout point à un instant donné. On dit qu'on travaille alors dans l'*approximation des régimes quasi-stationnaires* encore notée *A.R.Q.S.*. On parle également de régimes quasi-permanents au lieu de régimes quasi-stationnaires.

Dans la suite, on considère que les conditions de cette approximation sont réalisées. Sa justification sera traitée dans le cours de deuxième année : on verra que pour des circuits de taille raisonnable c'est-à-dire de longueur inférieure à un mètre et des fréquences inférieures à quelques mégahertz, les conditions de l'A.R.Q.S. sont obtenues. En effet, $\tau = \frac{L}{c} = \frac{1}{3 \cdot 10^8} \simeq 10^{-8}$ s et $T = \frac{1}{f} \leq 10^{-6}$ s. On a donc bien $\tau \ll T$, condition permettant de négliger les phénomènes de propagation.

4.4 Loi des nœuds

Dans l'A.R.Q.S., l'intensité est la même en tout point d'une branche du circuit. Cela signifie qu'il n'y a **pas d'accumulation de charges** en un point du circuit. Ainsi, pendant un intervalle de temps dt , il repart autant de charges d'un nœud qu'il en arrive.

Soit N un nœud où se rejoignent trois conducteurs (Cf. figure 1.6).

On note avec des majuscules les vraies intensités et avec des minuscules les intensités algébriques choisies pour étudier le phénomène.

Pendant l'intervalle de temps dt , il arrive la charge dq_a en N :

$$dq_a = (I_1 + I_2) dt$$

et il repart dq_r , telle que :

$$dq_r = I_3 dt$$

Ainsi puisqu'il n'y a pas d'accumulation de charges :

$$dq_a = dq_r \quad \Rightarrow \quad I_1 + I_2 = I_3 \quad (1.1)$$

Or avec les définitions des différentes intensités (Cf. figure 1.6), on peut écrire :

$$i_1 = I_1, \quad i_2 = -I_2, \quad i_3 = -I_3$$

En reportant ces relations dans l'expression (1.1), on trouve :

$$i_1 - i_2 + i_3 = 0 \quad (1.2)$$

Cette relation, qui porte le nom de *loi des nœuds*, se généralise au cas où n branches arrivent sur le nœud N , sous la forme suivante :

$$\sum_{k=1}^n \epsilon_k i_k = 0 \quad (1.3)$$

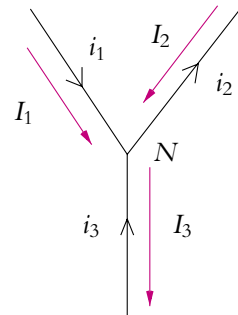


Figure 1.6 Nœud à trois branches.

avec $\epsilon_k = 1$ si l'intensité i_k est orientée vers le nœud N et $\epsilon_k = -1$ si l'intensité i_k est orientée depuis le nœud N .

4.5 Loi des mailles

Soit la maille de la figure 1.7.

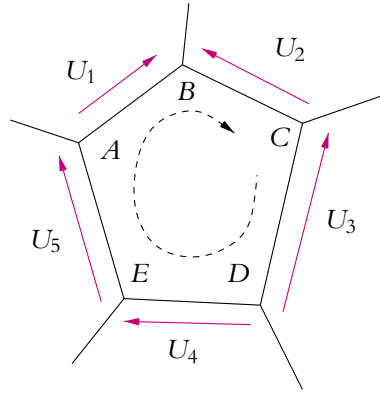


Figure 1.7 Maille à cinq branches.

On choisit un sens positif représenté par la flèche en pointillés. Sachant que la différence de potentiel entre le point A et lui-même est nulle, on peut écrire en introduisant les potentiels des autres points de la maille :

$$V_A - V_A = 0 \quad \Rightarrow \quad V_A - V_E + V_E - V_D + V_D - V_C + V_C - V_B + V_B - V_A = 0$$

Si on remplace les différences de potentiel par les tensions, avec :

$$U_1 = V_B - V_A, \quad U_2 = V_B - V_C, \quad U_3 = V_C - V_D, \quad U_4 = V_E - V_D \quad \text{et} \quad U_5 = V_A - V_E$$

on obtient la loi suivante :

$$U_1 - U_2 - U_3 + U_4 + U_5 = 0 \quad (1.4)$$

Cette loi, qui porte le nom de *loi des mailles*, se généralise à n tensions (n branches sur une maille) :

$$\sum_{k=1}^n \epsilon_k U_k = 0 \quad (1.5)$$

avec $\epsilon_k = 1$ si la tension U_k est dans le sens positif choisi et $\epsilon_k = -1$ si la tension U_k est dans le sens opposé au sens positif choisi.

5. Notion de dipôles

On appelle *dipôle électrocinétique* ou tout simplement *dipôle* tout système relié à un circuit électrique extérieur par deux bornes.

Quand on insère ce dipôle dans un circuit, une intensité électrique traverse en général ce dipôle et une tension s'installe à ses bornes. On précisera ultérieurement les conditions permettant un tel fonctionnement du dipôle.

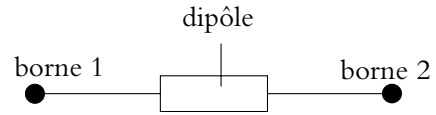


Figure 1.8 Définition d'un dipôle.

6. Puissance – dipôles récepteurs et générateurs

6.1 Définition de la puissance

Soit un dipôle parcouru par un courant d'intensité $i(t)$ et aux bornes duquel on a une tension $u(t) = V_A - V_B$.

On notera que du fait de la non accumulation de charges, l'intensité du courant entrant dans le dipôle et celle du courant sortant du dipôle sont les mêmes.

La puissance instantanée est par définition⁴ la quantité :

$$\mathcal{P}(t) = u(t) i(t)$$

Si on est en régime continu alors intensité et tension ne dépendent pas du temps et on peut écrire :

$$\mathcal{P} = u i$$

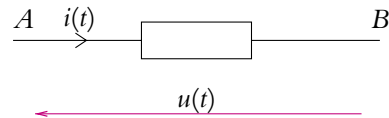


Figure 1.9 Convention pour la définition de la puissance.

6.2 Récepteurs et générateurs

Dans ce qui précède, on a décidé arbitrairement que le courant allait conventionnellement de A vers B (ce qui revient à dire dans la démonstration donnée en notes que A est l'entrée et B la sortie). D'autre part, on a défini la tension par :

$$u = V_A - V_B$$

⁴ Cette expression peut se justifier de la manière suivante en utilisant des résultats qui seront vus ultérieurement. Entre t et $t+dt$, la charge entrant dans le dipôle est $dq = idt$. On verra dans le cours d'électrostatique que l'énergie apportée ou cédée en un point s'écrit Vdq où V est le potentiel en ce point. Si on suppose que A est « l'entrée » du dipôle et B sa « sortie » alors l'énergie apportée est $V_A dq$ et l'énergie cédée est $V_B dq$, soit une variation d'énergie du dipôle : $dE = V_A dq - V_B dq = u dq$, ou encore $dE = u idt$. La puissance étant définie comme l'énergie reçue par unité de temps, on en déduit l'expression de la puissance.

soit une orientation de la tension opposée à celle de l'intensité.

Or il convient de noter les deux points importants suivants.

- L'intensité et la tension sont des grandeurs algébriques, elles peuvent être positives ou négatives suivant que l'orientation effective correspond ou non à l'orientation conventionnelle choisie, celle-ci étant choisie arbitrairement.
- Il existe deux possibilités d'orientations relatives de la tension et de l'intensité : de même sens ou de sens opposé.

Ces deux orientations relatives conduisent à deux conventions possibles :

- la *convention récepteur* où l'intensité i et la tension u sont choisies de sens opposé (c'est celle qu'on a adoptée au paragraphe précédent),

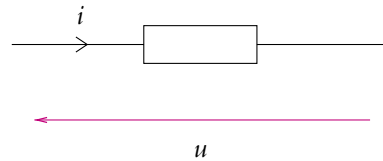


Figure 1.10 Convention récepteur.

- la *convention générateur* où l'intensité i et la tension u' sont choisies de même sens.

La définition de la puissance permet de donner une signification aux termes générateur et récepteur utilisés ici.

En convention récepteur, on a défini la puissance reçue par le dipôle par :

$$\mathcal{P}(t) = u(t) i(t)$$

Dans ce cas, si $u(t)$ et $i(t)$ sont positifs (cela correspond à l'adéquation du sens réel du courant et de la tension avec le sens conventionnel choisi), la puissance reçue est positive. Le dipôle est donc bien un récepteur : ce qu'on définit comme reçu l'est effectivement au sens où la quantité est positive.

A contrario, si $u(t)$ et $i(t)$ ne sont pas de même signe (ce qui signifie que l'une ou l'autre des quantités n'a pas une orientation conforme à la convention choisie) alors la puissance reçue définie dans le cadre de la convention récepteur est négative. Le dipôle se comporte comme un générateur qui fournit de la puissance au circuit.

Passer en convention générateur revient, par exemple, à inverser le sens conventionnel pour la tension et à considérer la tension $u'(t)$ telle que $u'(t) = -u(t)$. Dans ce cas, la puissance

$$\mathcal{P}_c(t) = u'(t)i(t)$$

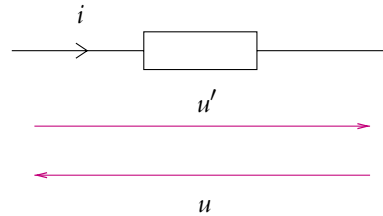


Figure 1.11 Convention générateur.

correspond à l'opposé de la puissance reçue utilisée précédemment. Il s'agit donc d'une puissance cédée qui sera effectivement positive lorsqu'il y aura diminution de l'énergie.

On peut résumer ces résultats par le tableau suivant :

Convention récepteur	Convention générateur
puissance reçue $\mathcal{P}(t) = u(t) i(t)$	puissance reçue $\mathcal{P}(t) = -u'(t) i(t)$
puissance cédée $\mathcal{P}_c(t) = -u(t) i(t)$	puissance cédée $\mathcal{P}_c(t) = u'(t) i(t)$

Lorsqu'on parle de puissance ou d'énergie, il est donc absolument nécessaire de :

- préciser la convention utilisée et donc l'orientation relative des tensions et des courants,
- préciser son caractère reçu ou cédé.

L'unité de la puissance est le watt, de symbole W, homogène à des joules par secondes. Il est à noter qu'au lieu d'utiliser le joule pour mesurer la consommation électrique, E.D.F. établit ses factures en kilowatt.heure, de symbole kW.h, ce qui correspond à une quantité d'énergie. L'équivalent en joules d'un kilowatt.heure est $1 \text{ kW.h} = 1\,000 \times 3\,600 = 3,6 \text{ MJ}$.

2

Circuits linéaires dans l'approximation des régimes quasi-stationnaires

Dans ce chapitre, on étudie les circuits linéaires c'est-à-dire ne comportant que des dipôles linéaires. On explicite la modélisation de quelques dipôles usuels en précisant leurs règles d'association ainsi que leur propriétés énergétiques. Quelques outils permettant de déterminer les intensités et les tensions d'un circuit sont également précisés.

1. Dipôles linéaires

On dit qu'un dipôle est *linéaire* si la tension à ses bornes $u(t)$ et l'intensité qui le traverse $i(t)$ sont liées par une équation différentielle linéaire à coefficients constants.

Le cas le plus simple est une relation affine entre l'intensité i et la tension u :

$$u = \alpha i + \beta$$

On verra que c'est le cas des résistors de résistance R .

Si l'intensité et/ou la tension varie en fonction des dérivées de l'une ou l'autre de ces grandeurs, il faut une équation différentielle qui est souvent du premier ordre :

$$a_1 \frac{du}{dt} + a_0 u + b_1 \frac{di}{dt} + b_0 i = f(t)$$

en notant $f(t)$ une fonction du temps indépendante de la tension u et de l'intensité i et a_1 , a_0 , b_1 et b_0 des constantes. On verra au cours des paragraphes suivants que c'est le cas des bobines d'inductance L et des condensateurs de capacité C .

On peut généraliser cette définition sous la forme suivante :

$$\sum_{k=0}^K a_k \frac{d^k u}{dt^k} + \sum_{l=1}^L b_l \frac{d^l i}{dt^l} = f(t)$$

avec $f(t)$ indépendant de u et de i et $\forall (k, l)$, a_k et b_l constants.

2. Résistor de résistance R

2.1 Représentation

Ce dipôle est schématisé en convention récepteur par :

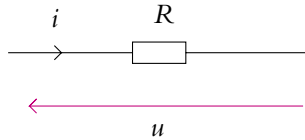


Figure 2.1 Symbole d'une résistance.

2.2 Caractéristique

Il s'agit du dipôle qui vérifie la loi d'Ohm en convention récepteur :

$$u = Ri$$

R est appelé *résistance*, elle est positive et s'exprime en ohms, de symbole Ω . On peut également définir la *conductance* G comme l'inverse de la résistance :

$$G = \frac{1}{R}$$

G s'exprime en Ω^{-1} ou en siemens, de symbole S . En convention récepteur, la loi d'Ohm s'écrit aussi :

$$i = Gu$$

On peut représenter cette relation en traçant l'intensité i traversant le résistor en fonction de la tension à ses bornes (on dit qu'on trace la caractéristique courant-tension du résistor) :

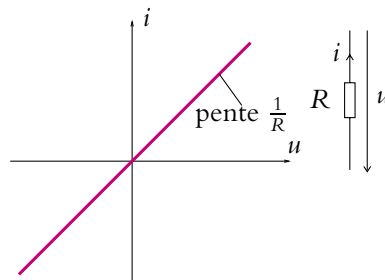


Figure 2.2 Caractéristique courant-tension d'un résistor en convention récepteur.

2.3 Association en série

Cette association consiste à placer les dipôles de telle sorte que la même intensité traverse les dipôles :

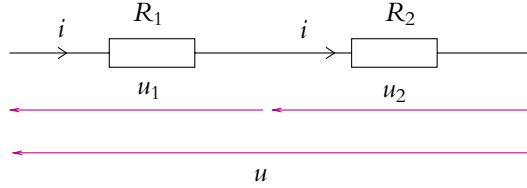


Figure 2.3 Association en série de deux résistors.

On en déduit que la tension aux bornes de l'ensemble est la somme des tensions aux bornes de chaque dipôle :

$$u = u_1 + u_2$$

On peut généraliser ce résultat au cas de N dipôles :

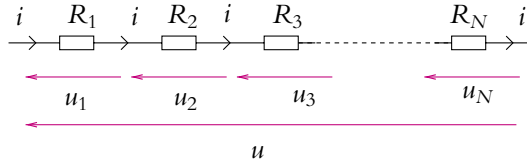


Figure 2.4 Association en série de résistors.

Les N dipôles sont en série si une même intensité traverse tous les dipôles :

$$i_1 = i_2 = \dots = i_N = i$$

La tension aux bornes de l'ensemble est la somme des tensions aux bornes de chaque dipôle :

$$u = u_1 + u_2 + \dots + u_N$$

Dans le cas où les dipôles sont des résistors de résistance $R_1, R_2, \dots, R_N$:

$$u = R_1 i + R_2 i + \dots + R_N i = (R_1 + R_2 + \dots + R_N) i$$

L'association en série de résistors de résistance $R_1, R_2, \dots, R_N$ est donc un résistor de résistance

$$R = R_1 + R_2 + \dots + R_N$$

ou de conductance G telle que :

$$\frac{1}{G} = \frac{1}{G_1} + \frac{1}{G_2} + \dots + \frac{1}{G_N}$$

2.4 Association en parallèle

Cette association correspond au cas où les deux dipôles ont même tension à leurs bornes, selon le schéma suivant :

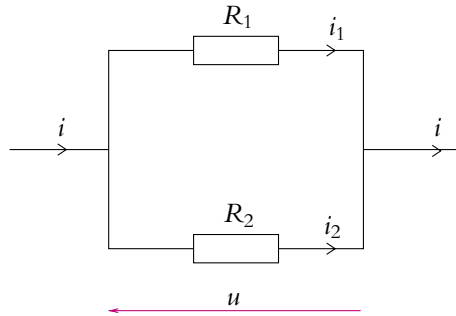


Figure 2.5 Association en parallèle de deux résistors.

On en déduit que l'intensité entrant ou sortant de l'association parallèle est la somme des intensités traversant chaque dipôle :

$$i = i_1 + i_2$$

On peut généraliser ce résultat au cas de N dipôles :

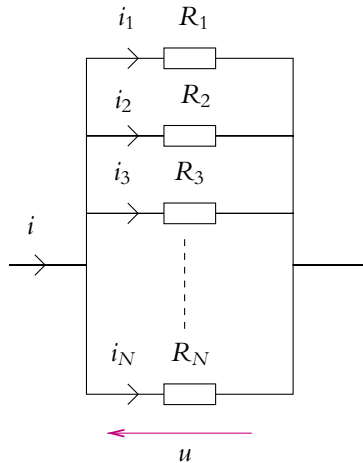


Figure 2.6 Association en parallèle de résistors.

Les N dipôles sont en parallèle si la tension aux bornes de tous les dipôles est la même :

$$u_1 = u_2 = \dots = u_N = u$$

L'intensité entrant ou sortant de l'association parallèle est la somme des intensités traversant chaque dipôle :

$$i = i_1 + i_2 + \dots + i_N$$

On notera que ce dernier résultat est tout simplement la loi des nœuds.

Dans le cas où les dipôles sont des résistors de conductance $G_1, G_2, \dots, G_N$:

$$i = G_1 u + G_2 u + \dots + G_N u = (G_1 + G_2 + \dots + G_N) u$$

L'association en parallèle de résistors de conductance $G_1, G_2, \dots, G_N$ est donc un résistor de conductance

$$G = G_1 + G_2 + \dots + G_N$$

ou de résistance R telle que :

$$\frac{1}{R} = \frac{1}{R_1} + \frac{1}{R_2} + \dots + \frac{1}{R_N}$$

2.5 Puissance dissipée dans un résistor

En convention récepteur, la loi d'Ohm s'écrit $u(t) = Ri(t)$, la puissance reçue peut donc se mettre sous la forme :

$$\mathcal{P} = Ri^2 \quad \text{ou} \quad \mathcal{P} = \frac{u^2}{R} \quad (2.1)$$



La puissance reçue par un résistor est toujours positive : un résistor se comporte toujours en récepteur.

L'énergie reçue entre les instants t et $t + dt$ est donc : $\delta W = Ri^2 dt = \frac{u^2}{R} dt$ et entre t_1 et t_2 :

$$W = \int_{t_1}^{t_2} Ri^2(t) dt = \int_{t_1}^{t_2} \frac{u^2(t)}{R} dt$$

Pour terminer le calcul, il faudrait connaître les expressions de $i(t)$ et $u(t)$.

En pratique, cette énergie est dissipée sous forme de transfert thermique : il s'agit de l'effet Joule.

3. Bobine d'inductance L

3.1 Bobine et auto-induction

Une bobine est constituée d'un enroulement de spires conductrices.

On reverra dans le cours de deuxième année que le phénomène dit d'auto-induction crée aux bornes d'une bobine une tension u lorsque le courant d'intensité i qui la traverse varie au cours du temps. La traduction mathématique de ce phénomène est la relation suivante entre u et i en convention récepteur :

$$u = L \frac{di}{dt}$$

L est appelée *inductance* et s'exprime en henry, de symbole H. On la représente en convention récepteur par :

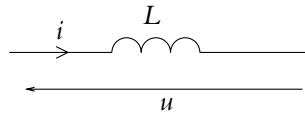


Figure 2.7 Symbole d'une inductance.

Si on utilise la convention générateur, on a alors la relation

$$u' = -u = -L \frac{di}{dt}$$

Ce phénomène a déjà été vu en termes de force électromotrice dans les classes antérieures. Ici le seul point qui importe est la relation entre intensité et tension ; elle sera admise et on n'étudiera pas davantage le phénomène.



En régime continu, i est une constante et la relation précédente implique que $u = 0$: la bobine constitue alors un court-circuit.

On verra dans le chapitre sur l'instrumentation électrique que la modélisation d'une bobine réelle nécessite de tenir compte d'une résistance interne due aux enroulements.

3.2 Énergie emmagasinée dans une bobine d'inductance L

En convention récepteur, la relation tension – courant s'écrit pour une bobine d'inductance L : $u = L \frac{di}{dt}$. La puissance reçue se met alors sous la forme :

$$\mathcal{P} = u(t)i(t) = L \frac{di}{dt} i(t) = \frac{d}{dt} \left(\frac{1}{2} L i^2 \right)$$

Sachant que la puissance est la dérivée de l'énergie par rapport au temps, l'expression précédente fait apparaître l'énergie instantanée d'une bobine d'inductance L , c'est-à-dire présente dans la bobine à un instant donné :

$$E = \frac{1}{2}Li^2.$$

L'énergie reçue entre deux instants t_1 et t_2 est donc :

$$W = \int_{t_1}^{t_2} \frac{d}{dt} \left(\frac{1}{2}Li^2 \right) dt = \left[\frac{1}{2}Li^2 \right]_{t_1}^{t_2} = \frac{1}{2}Li^2(t_2) - \frac{1}{2}Li^2(t_1) = E(t_2) - E(t_1)$$

L'énergie est une fonction continue du temps c'est-à-dire qu'elle ne peut pas apparaître subitement. On déduit de la relation précédente que **l'intensité parcourant une bobine est une fonction continue du temps**.

La puissance reçue par une bobine peut changer de signe au cours du temps.

Si E diminue (donc si $|i|$ diminue), $\mathcal{P}$ est négative : la bobine cède effectivement de l'énergie à l'extérieur et se comporte comme un générateur.

En revanche, si E augmente, $\mathcal{P}$ est positive : la bobine reçoit effectivement de l'énergie de l'extérieur et se comporte comme un récepteur.

3.3 Complément : association en série

On a vu que l'association en série de dipôles vérifiait :

$$u = u_1 + u_2 + \dots + u_N,$$

les dipôles étant parcourus par la même intensité i . Pour le cas où les dipôles sont des bobines d'inductances $L_1, L_2, \dots, L_N$:

$$u = L_1 \frac{di}{dt} + L_2 \frac{di}{dt} + \dots + L_N \frac{di}{dt} = (L_1 + L_2 + \dots + L_N) \frac{di}{dt}$$

L'association en série de bobines d'inductances $L_1, L_2, \dots, L_N$ est donc une bobine d'inductance :

$$L = L_1 + L_2 + \dots + L_N$$

On retrouve la même loi d'association que pour les résistances.

3.4 Complément : association en parallèle

On a vu que l'association en parallèle vérifiait :

$$i = i_1 + i_2 + \dots + i_N,$$

les dipôles ayant la même tension à leurs bornes. Pour le cas où les dipôles sont des bobines d'inductances $L_1, L_2, \dots, L_N$:

$$u = L_1 \frac{di_1}{dt} = L_2 \frac{di_2}{dt} = \dots = L_N \frac{di_N}{dt}$$

or $i = i_1 + i_2 + \dots + i_N$ donc :

$$\frac{di}{dt} = \frac{di_1}{dt} + \frac{di_2}{dt} + \dots + \frac{di_N}{dt}$$

soit :

$$\begin{aligned} \frac{di}{dt} &= \frac{u}{L_1} + \frac{u}{L_2} + \dots + \frac{u}{L_N} \\ &= \left(\frac{1}{L_1} + \frac{1}{L_2} + \dots + \frac{1}{L_N} \right) u \end{aligned}$$

L'association en parallèle de bobines d'inductances est donc une bobine inductance L telle que :

$$\frac{1}{L} = \frac{1}{L_1} + \frac{1}{L_2} + \dots + \frac{1}{L_N}$$

On retrouve les mêmes lois d'association en parallèle que celles obtenues pour des résistances.

4. Condensateur de capacité C

4.1 Condensateur

Les condensateurs sont des composants constitués de :

- deux conducteurs qui se font face et sont appelés armatures,
- un matériau isolant, le diélectrique, situé entre les deux armatures.

Ils peuvent être de plusieurs formes : plan, cylindrique, etc.

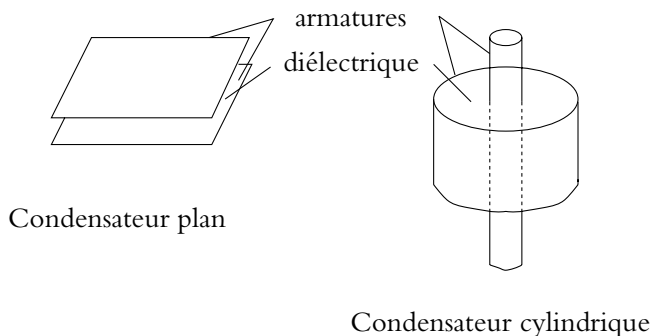


Figure 2.8 Structure d'un condensateur.

En électricité, on utilise la plupart du temps des condensateurs plans enroulés pour des raisons de gain de place¹.

L'une des armatures porte une charge q tandis que l'autre porte une charge $-q$. La modélisation la plus simple des condensateurs est celle d'une capacité C .

4.2 Définition et représentation d'une capacité

Une *capacité* C est caractérisée par la relation² entre la charge q et la tension appliquée aux bornes u :

$$q = Cu$$

On la représente par :

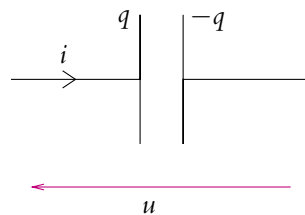


Figure 2.9 Symbole d'une capacité.

On notera l'importance du sens choisi pour l'intensité i par rapport à la position des charges q et $-q$: l'intensité i arrive sur l'armature de charge $+q$.

Les capacités sont exprimées en farads, de symbole F.

4.3 Relation tension - intensité

L'arrivée d'un courant i sur une armature provoque une variation dq de la charge de l'armature et donc une variation $-dq$ sur l'autre pour assurer la conservation de la charge au niveau du composant. Un courant d'intensité i partira de la seconde armature même si les charges ne traversent pas physiquement l'isolant. On obtient bien un dipôle au sens où cela a été introduit.

Pendant l'intervalle de temps dt , il arrive une charge $\delta q = idt$ sur l'armature de charge $+q$ et il repart $\delta q' = -idt$ sur l'armature de charge $-q$. Or la conservation de la charge de la première armature se traduit par :

$$q(t + dt) = q(t) + \delta q = q(t) + idt$$

¹ Il existe des condensateurs chimiques d'utilisation différente : ils ne fonctionnent que dans un seul sens, ils sont dits polarisés.

² Cette relation sera établie dans le cours d'électromagnétisme.

Or, au premier ordre en dt , $q(t + dt) = q(t) + \frac{dq}{dt}dt$. On en déduit :

$$i = \frac{dq}{dt}$$

où q est la charge de l'armature qui voit arriver le courant d'intensité i .

Comme $q = Cu$, on peut en déduire la relation

$$i = C \frac{du}{dt}$$

entre l'intensité i parcourant le condensateur et la tension u à ses bornes.



En régime continu, la tension aux bornes du condensateur est constante et l'intensité du courant est donc nulle : $i = 0$. Par conséquent, un condensateur se comporte en régime continu comme un interrupteur ouvert.

4.4 Énergie emmagasinée dans un condensateur de capacité C

En convention récepteur, la relation tension - courant s'écrit pour un condensateur de capacité C : $i = C \frac{du}{dt}$. La puissance reçue par le condensateur se met sous la forme :

$$\mathcal{P} = u(t)i(t) = u(t)C \frac{du}{dt} = \frac{d}{dt} \left(\frac{1}{2} Cu^2 \right)$$

Comme dans le cas de la bobine, l'expression précédente fait apparaître l'énergie instantanée d'un condensateur de capacité C , c'est-à-dire présente dans le condensateur à un instant donné :

$$E = \frac{1}{2} Cu^2.$$

L'énergie reçue entre deux instants t_1 et t_2 est donc :

$$W = \int_{t_1}^{t_2} \frac{d}{dt} \left(\frac{1}{2} Cu^2 \right) dt = \left[\frac{1}{2} Cu^2 \right]_{t_1}^{t_2} = \frac{1}{2} Cu^2(t_2) - \frac{1}{2} Cu^2(t_1) = E(t_2) - E(t_1)$$

L'énergie est une grandeur continue dans le temps. De l'expression de l'énergie instantanée, on déduit que **la tension aux bornes d'un condensateur de capacité C est une fonction continue du temps ainsi que la charge qui lui est proportionnelle**. Il s'agit du même raisonnement que celui qui a permis d'établir que l'intensité traversant une bobine d'inductance L est continue.

La puissance reçue par un condensateur peut changer de signe au cours du temps. Si E diminue (donc si $|u|$ diminue), $\mathcal{P}$ est négative : le condensateur cède effectivement de l'énergie à l'extérieur et se comporte comme un générateur. En revanche,

si E augmente, $\mathcal{P}$ est positive : le condensateur reçoit effectivement de l'énergie de l'extérieur et se comporte comme un récepteur.

4.5 Complément : association en série

On a vu que l'association en série vérifiait :

$$u = u_1 + u_2 + \dots + u_N,$$

les dipôles étant parcourus par le même courant. Pour le cas où les dipôles sont des condensateurs de capacités C_1 et C_2 :

$$\left\{ \begin{array}{l} i = C_1 \frac{du_1}{dt} \\ i = C_2 \frac{du_2}{dt} \\ \vdots \\ i = C_N \frac{du_N}{dt} \end{array} \right.$$

Or $u = u_1 + u_2 + \dots + u_N$, donc :

$$\frac{du}{dt} = \frac{du_1}{dt} + \frac{du_2}{dt} + \dots + \frac{du_N}{dt}$$

soit :

$$\frac{du}{dt} = \frac{i}{C_1} + \frac{i}{C_2} + \dots + \frac{i}{C_N} = \left(\frac{1}{C_1} + \frac{1}{C_2} + \dots + \frac{1}{C_N} \right) i$$

L'association en série de condensateurs de capacités $C_1, C_2, \dots, C_N$ est un condensateur de capacité C telle que :

$$\frac{1}{C} = \frac{1}{C_1} + \frac{1}{C_2} + \dots + \frac{1}{C_N}$$

L'association en série de capacités est donc analogue à celle des conductances.

4.6 Complément : association en parallèle

On a vu que l'association en parallèle vérifiait :

$$i = i_1 + i_2 + \dots + i_N,$$

les dipôles ayant même tension à leurs bornes. Pour le cas où les dipôles sont des condensateurs de capacités $C_1, C_2, \dots, C_N$:

$$i = C_1 \frac{du}{dt} + C_2 \frac{du}{dt} + \dots + C_N \frac{du}{dt} = (C_1 + C_2 + \dots + C_N) \frac{du}{dt}$$

L'association en parallèle de condensateurs de capacités $C_1, C_2, \dots, C_N$ est donc un condensateur de capacité C :

$$C = C_1 + C_2 + \dots + C_N$$

Elle est donc analogue à l'association en parallèles des conductances.

5. Sources de tension et de courant - Modèles de Thévenin et de Norton

5.1 Source de tension

On appelle *source de tension* un dispositif idéal qui impose une différence de potentiel constante aux bornes du circuit auquel il est relié, quelle que soit l'intensité du courant qui le traverse.

Sa représentation en convention générateur et sa caractéristique sont les suivantes :

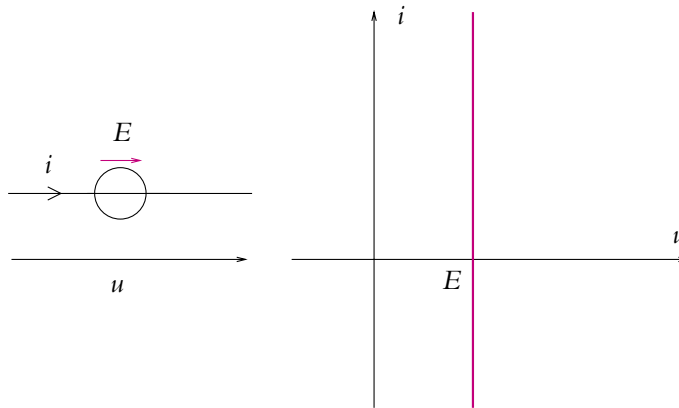


Figure 2.10 Symbole et caractéristique d'une source idéale de tension.

En effet, la tension est indépendante de l'intensité du courant parcourant le circuit par définition même du composant. On parle également de *force électromotrice* de la source, soit f.e.m. en abrégé, ou de *tension à vide* (la tension étant la même pour tout courant, c'est notamment celle correspondant au cas $i = 0$ qui est la tension à vide par définition).

► **Remarque** : Il faudra donc porter une attention particulière à ce type de dipôles car si la tension à ses bornes est connue, il n'en est rien de l'intensité qui le traverse : elle peut *a priori* prendre toutes les valeurs possibles.

5.2 Sources de courant

On appelle *source de courant* un dispositif idéal qui débite un courant d'intensité constante dans le circuit auquel il est relié quelle que soit la tension à ses bornes et ce indépendamment du circuit³.

Son symbole en convention générateur et sa caractéristique sont représentés sur la figure 2.11.

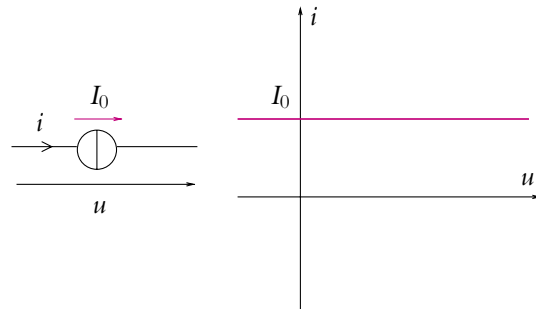


Figure 2.11 Symbole et caractéristique d'une source idéale de courant.

La grandeur I_0 est appelée *courant de court-circuit* : elle est indépendante du circuit qui peut être n'importe quel dispositif et en particulier un fil de connexion créant un court-circuit. On utilise aussi l'expression courant électromoteur ou c.e.m. par analogie aux sources de tension.

► **Remarque** : la valeur de l'intensité est indépendante de la valeur de la tension : la donnée de l'intensité ne fixe pas celle de la tension. La tension peut prendre n'importe quelle valeur.

5.3 Modèle de Thévenin

Le plus souvent, la caractéristique tension-courant des générateurs a l'allure suivante (on a choisi ici de représenter u en fonction de i pour que la pente de la droite soit homogène en valeur absolue à celle d'une résistance) :

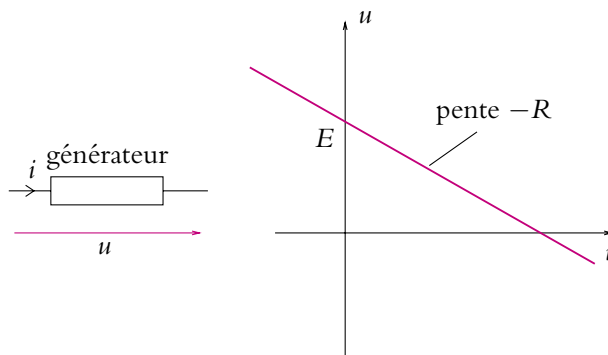


Figure 2.12 Caractéristique générale d'un dipôle linéaire.

³ Il s'agit du « dual » pour le courant des sources de tension. On aura des résultats analogues en inversant tension et courant, on parle alors de dualité.

Il s'agit de la caractéristique d'un dipôle linéaire dont l'équation peut s'écrire :

$$u = E - Ri \quad (2.2)$$

On peut alors modéliser ce dipôle par une source de tension idéale et une résistance en série :

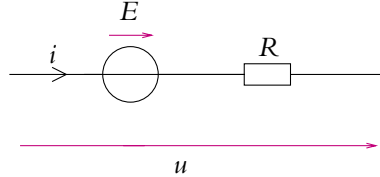


Figure 2.13 Modèle de Thévenin d'un dipôle linéaire.

En effet, en convention générateur, la tension u aux bornes du dipôle et la f.e.m. E sont de même sens et la relation entre intensité et tension pour une résistance en convention générateur est $u = -Ri$ soit par association en série $u = E - Ri$.

D'autres choix sont possibles, par exemple :

$$u = -E - Ri$$

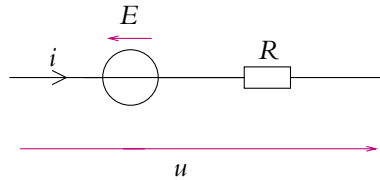


Figure 2.14 Autre modèle de Thévenin d'un dipôle linéaire.

ou

$$u = E + Ri$$

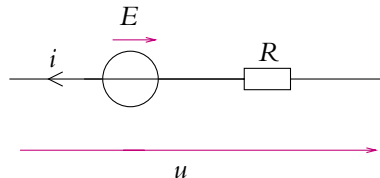


Figure 2.15 Autre modèle de Thévenin d'un dipôle linéaire.

Il s'agit du modèle dit *modèle de Thévenin*. On appelle force électromotrice du générateur la grandeur E et résistance interne la grandeur R .

5.4 Modèle de Norton

La caractéristique précédente est également la courbe d'équation

$$i = \frac{E}{R} - \frac{1}{R}u \quad (2.3)$$

Cela correspond à la relation :

$$i = I_0 - Gu$$

avec $I_0 = \frac{E}{R}$ et $G = \frac{1}{R}$.

Il est donc possible de considérer une deuxième modélisation du générateur :

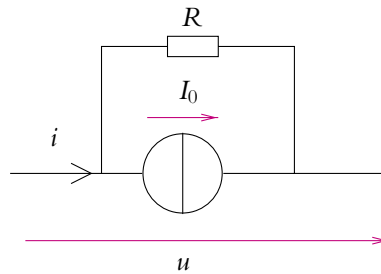


Figure 2.16 Modèle de Norton d'un dipôle linéaire.

Il s'agit du modèle dit *modèle de Norton* qui est l'association en parallèle d'une source idéale de courant, de courant de court-circuit I_0 , et d'une résistance R avec

$$I_0 = \frac{E}{R}$$

5.5 Transformation Thévenin - Norton

On a modélisé dans les deux paragraphes précédents un même dipôle linéaire en exploitant différemment l'équation de sa caractéristique. On obtient ainsi l'équivalence des deux modélisations de Thévenin et de Norton :

Modèle de Norton	Modèle de Thévenin
conductance G	résistance R
courant de court-circuit $I_0 = \frac{E}{R}$	force électromotrice $E = RI_0$
$i = I_0 - \frac{1}{R}u = I_0 - Gu$	$u = E - Ri$

Les représentations de Thévenin et de Norton d'un dipôle linéaire ainsi que le passage de l'une à l'autre sont basés sur la linéarité de la relation entre la tension aux bornes du dipôle et l'intensité du courant qui le traverse : elles s'appliquent donc à tout dipôle linéaire, c'est-à-dire à tout dipôle dont la relation entre u et i est de la forme (2.2) ou (2.3).

6. Lois de Kirchhoff

On englobe sous le nom de *lois de Kirchhoff* deux types d'expressions permettant de calculer soit les intensités dans toutes les branches du circuit considéré, soit toutes les tensions entre les nœuds du circuit. La première méthode est appelée *méthode des mailles*, la seconde *méthode des nœuds*.

Dans la suite du paragraphe, on développe ces méthodes sur des exemples simples avant de donner une procédure plus générale.

Ces méthodes deviennent vite fastidieuses au fur et à mesure que le nombre de mailles augmente. **On limitera leur utilisation aux cas de circuits à faible nombre de mailles.**

6.1 Méthode des mailles

Les inconnues recherchées sont les intensités dans toutes les branches du circuit. Dans ce cas, il faut transformer tous les générateurs en modèle de Thévenin puis appliquer des lois des mailles.

On considère le circuit de la figure 2.17, les inconnues sont les intensités i_1 , i_2 et i_3 qu'on cherche à exprimer en fonction des résistances R , R_1 et R_2 ainsi que de la force électromotrice E et du courant de court-circuit I .

La première étape consiste à transformer le modèle de Norton du générateur en modèle de Thévenin équivalent, ce qui donne le schéma de la figure 2.18.

On écrit ensuite les lois des mailles relatives aux mailles (1) et (2) en prenant les sens positifs indiqués sur la figure 2.18 :

$$\begin{cases} E' - V_1 - V_3 = 0 \\ V_3 - V_2 - E = 0 \end{cases}$$

où $E' = R_1 I$

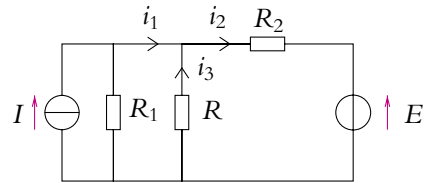


Figure 2.17 Circuit initial.

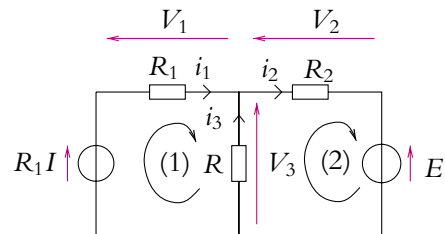


Figure 2.18 Circuit après transformation du générateur de courant.

► **Remarque :** Le circuit comporte en fait trois mailles : (1), (2) et la maille globale contenant E' , E , R_1 et R_2 . Si on écrit la loi des mailles relative à cette dernière, on peut facilement se convaincre que l'équation obtenue n'apporte pas d'information supplémentaire : il s'agit de la somme des deux relations précédentes. On n'a donc que deux mailles indépendantes sur les trois.

La loi d'Ohm appliquée aux résistors donne les relations suivantes :

$$\begin{cases} V_1 = R_1 i_1 \\ V_2 = R_2 i_2 \\ V_3 = -R i_3 \end{cases}$$

Le signe « $-$ » de la dernière relation est dû au fait que R est en convention générateur. Finalement on doit résoudre le système :

$$\begin{cases} R_1 i_1 - R i_3 = E' \\ R_2 i_2 + R i_3 = -E \end{cases} \quad (2.4)$$

On ne dispose alors que de deux équations pour trois inconnues : il manque une relation qui sera obtenue en écrivant la loi des nœuds :

$$i_3 = i_2 - i_1 \quad (2.5)$$

Finalement, en reportant l'expression de i_3 dans (2.4), on obtient :

$$\begin{cases} (R_1 + R) i_1 - R i_2 = E' \\ -R i_1 + (R_2 + R) i_2 = -E \end{cases} \quad (2.6)$$

La résolution de ce système donne :

$$\begin{cases} i_1 = \frac{-RE + (R_2 + R)E'}{R_1 R_2 + R_1 R + R_2 R} \\ i_2 = \frac{-(R_1 + R)E + RE'}{R_1 R_2 + R_1 R + R_2 R} \end{cases} \quad (2.7)$$



On ne doit jamais avoir de différence au dénominateur, sinon ce dernier pourrait s'annuler pour un choix *ad'hoc* des résistances : les intensités deviendraient alors infinies, ce qui n'a pas de sens physique.

On peut tirer une loi générale du système (2.6) obtenu. On remarque qu'on compte positivement les f.e.m. lorsqu'elles sont dans le sens positif choisi et négativement dans le cas contraire. De même, on compte positivement les courants quand ils sont dans le sens positif choisi et négativement dans le cas contraire. On peut résumer ceci

par l'équation suivante, pour une maille contenant N générateurs de f.e.m. E_j et L résistors de résistances R_k parcourues par les intensités i_k :

$$\sum_{j=1}^N \epsilon_j E_j = \sum_{k=1}^L R_k \mu_k i_k \quad (2.8)$$

avec :

- $\epsilon_j = 1$ si E_j est dans le sens positif choisi et $\epsilon_j = -1$ sinon,
- $\mu_k = 1$ si i_k est dans le sens positif choisi et $\mu_k = -1$ sinon.

Récapitulatif des différentes étapes :

- On transforme tous les générateurs en modèle de Thévenin et on représente le circuit avec toutes les orientations conventionnelles des intensités.
- On écrit une équation du type (2.8) pour toutes les mailles indépendantes. En effet, comme on l'a vu sur le circuit de la figure 2.18, si le circuit possède M mailles, seules $(M - 1)$ sont indépendantes. On notera qu'il faut que chaque branche intervienne dans au moins une équation de mailles et que, pour assurer l'indépendance des équations, il suffit d'avoir une branche dont on n'a pas encore tenu compte à chaque nouvelle équation des mailles écrite.
- On écrit les lois des nœuds reliant les intensités de manière à n'avoir plus que le nombre d'inconnues correspondant au nombre d'équations
- On résout le système d'équations.

6.2 Méthode des nœuds

Les inconnues sont les potentiels des nœuds ou les tensions par rapport à un nœud de référence (en général la masse du circuit). Dans ce cas, il faut transformer tous les générateurs indépendants en modèle de Norton puis appliquer des lois des nœuds.

On considère le circuit de la figure 2.19, les inconnues sont les tensions V_{12} , V_{13} et V_{23} qu'on cherche à exprimer en fonction des résistances R_1 , R_2 et R_3 ainsi que de la force électromotrice E et des courants de court-circuit I et J .

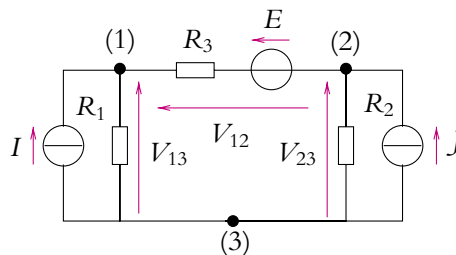


Figure 2.19 Circuit initial.

On remarque que seules deux de ces trois tensions sont indépendantes puisque $V_{12} = V_{13} - V_{23}$.

La première étape consiste à transformer le modèle de Thévenin du générateur en modèle de Norton équivalent, ce qui donne le schéma de la figure 2.20.

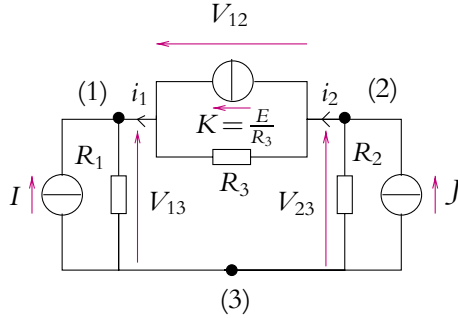


Figure 2.20 Circuit après transformation du générateur de tension.

On écrit ensuite les lois des nœuds aux nœuds (1) et (2) en tenant compte des conventions générateur et récepteur utilisées :

$$\begin{cases} i_1 = -I + G_1 V_{13} = K - G_3 V_{12} \\ i_2 = K - G_3 V_{12} = J - G_2 V_{23} \end{cases}$$

en notant G_1 , G_2 et G_3 les conductances correspondant aux résistances R_1 , R_2 et R_3 et $K = \frac{E}{R_3}$ (Cf. figure 2.21).

► **Remarque :** Le circuit comporte en fait 3 nœuds (1), (2) et (3). On peut facilement se convaincre que l'équation obtenue en écrivant la loi des nœuds au nœud (3) s'obtient en sommant les deux équations précédentes, elle n'apporte pas de renseignement supplémentaire.

Comme $V_{12} = V_{13} - V_{23}$, on en déduit finalement :

$$\begin{cases} I + K = (G_1 + G_3)V_{13} - G_3 V_{23} \\ J - K = -G_3 V_{13} + (G_2 + G_3)V_{23} \end{cases} \quad (2.9)$$

La résolution de système donne :

$$\begin{cases} V_{13} = \frac{(G_2 + G_3)I + G_3 J + G_2 K}{G_1 G_2 + G_1 G_3 + G_2 G_3} \\ V_{23} = \frac{G_3 I + (G_1 + G_3)J - G_1 K}{G_1 G_2 + G_1 G_3 + G_2 G_3} \end{cases} \quad (2.10)$$



Comme pour la méthode des mailles, on ne doit jamais avoir de différence au dénominateur, sinon ce dernier pourrait s'annuler pour un choix *ad'hoc* des conductances : les tensions deviendraient alors infinies, ce qui n'a pas de sens physique.

On peut tirer une loi générale du système (2.9) obtenu. On remarque qu'on compte positivement les c.e.m. lorsqu'ils sont dirigés vers le nœud considéré et négativement dans le cas contraire. De même, on compte positivement les tensions quand elles ont la pointe de la flèche vers le nœud choisi et négativement dans le cas contraire. On peut résumer ceci par l'équation suivante, pour un nœud auquel sont reliés N générateurs de c.e.m. I_j et L conductances G_k soumises aux tensions V_k :

$$\sum_{j=1}^N \epsilon_j I_j = \sum_{k=1}^L G_k \mu_k V_k \quad (2.11)$$

avec :

- $\epsilon_j = 1$ si I_j est vers le nœud et $\epsilon_j = -1$ sinon,
- $\mu_k = 1$ si V_k est vers le nœud et $\mu_k = -1$ sinon.

Récapitulatif des différentes étapes :

- On transforme tous les générateurs en modèle de Norton et on représente le circuit avec toutes les orientations conventionnelles des tensions et des intensités.
- On écrit une équation du type (2.11) pour tous les nœuds indépendants. En effet, comme on l'a vu sur l'exemple du circuit 2.20, si le circuit possède N nœuds, seuls $(N - 1)$ sont indépendants. On notera qu'il faut que chaque dipôle intervienne dans au moins une équation des nœuds et que, pour assurer l'indépendance des équations, il suffit d'avoir un dipôle dont on n'a pas encore tenu compte à chaque nouvelle équation des nœuds écrite.
- On écrit les lois des mailles reliant les tensions de manière à n'avoir plus que le nombre d'inconnues correspondant au nombre d'équations.
- On résout le système d'équations.

7. Diviseurs de tension et de courant

On s'intéresse dans ce paragraphe aux diviseurs de tension et de courant, deux types de circuits souvent présents dans les circuits.

7.1 Diviseur de tension

La structure de base du diviseur de tension est donnée par le montage de la figure 2.21 : on soumet l'association en série de deux résistors à une tension V_e et on cherche la tension aux bornes de l'un d'entre eux.

Les expressions des tensions en fonction du courant parcourant les résistors sont :

$$\begin{cases} V_1 = R_1 i \\ V_2 = R_2 i \\ V_e = (R_1 + R_2) i \end{cases}$$

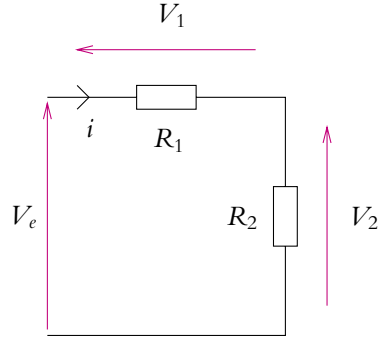


Figure 2.21 Diviseur de tension.

En faisant le rapport des expressions, on trouve :

$$V_1 = \frac{R_1}{R_1 + R_2} V_e \quad \text{et} \quad V_2 = \frac{R_2}{R_1 + R_2} V_e \quad (2.12)$$

Les tensions V_1 et V_2 sont des fractions de la tension totale V_e , ce qui explique la dénomination « diviseur de tension » donnée à ce circuit. Le rapport de la tension aux bornes d'une résistance à la tension totale est égal au rapport de la résistance considérée à la résistance totale.

Ce résultat est généralisable à plus de deux résistors en série : pour N résistors en série soumis à la tension totale V_e , la tension V_k aux bornes du résistor de résistance R_k est :

$$V_k = \frac{R_k}{\sum_{j=1}^N R_j} V_e \quad (2.13)$$

► Deux remarques importantes :

- Il faut faire attention à appliquer correctement la formule du diviseur de tension notamment quand on a des associations de résistances en parallèle dans le circuit. Par exemple, pour le circuit de la figure 2.22, le calcul de la tension V nécessite de tenir compte du fait que R_1 et R_2 ne sont pas en série : on n'aura pas $V = \frac{R_2 V_e}{R_1 + R_2}$. Il faut remplacer R_2 et R_3 par leur résistance équivalente $R_{eq} = \frac{R_2 R_3}{R_2 + R_3}$ pour pouvoir appliquer la relation du diviseur de tension à R_1 et R_{eq} :

$$V = \frac{R_{eq}}{R_1 + R_{eq}} V_e$$

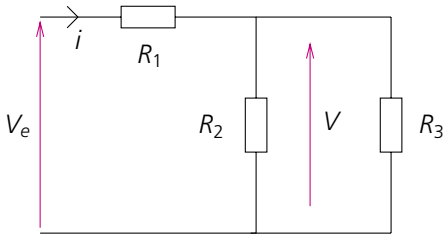


Figure 2.22 Diviseur de tension sur une charge.

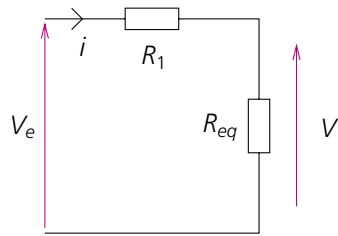


Figure 2.23 Schéma équivalent à un diviseur de tension sur une charge.

- Il faut également faire attention aux sens des orientations des tensions.

7.2 Diviseur de courant

La structure de base du diviseur de courant est donnée par le montage de la figure 2.24 : on soumet l'association en parallèle de deux résistors à un courant d'intensité totale i et on cherche l'intensité parcourant l'un d'entre eux.

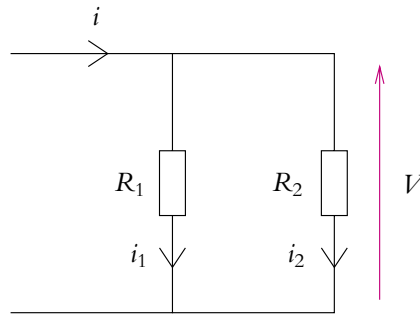


Figure 2.24 Diviseur de courant.

En utilisant les conductances $G_1 = \frac{1}{R_1}$ et $G_2 = \frac{1}{R_2}$, on peut écrire les relations suivantes :

$$\begin{cases} i_1 = G_1 V \\ i_2 = G_2 V \\ i = (G_1 + G_2) V \end{cases}$$

En faisant le rapport des expressions, on trouve :

$$i_1 = \frac{G_1}{G_1 + G_2} i \quad \text{et} \quad i_2 = \frac{G_2}{G_1 + G_2} i \quad (2.14)$$

Les intensités i_1 et i_2 sont des fractions de l'intensité totale i , ce qui explique la dénomination « diviseur de courant » donnée à ce circuit. Le rapport de l'intensité parcourant

une conductance à l'intensité totale est égal au rapport de la conductance considérée à la conductance totale.

Ce résultat est généralisable à plus de deux résistors en parallèle : pour N résistors en parallèle soumis à l'intensité totale i , l'intensité i_k dans le résistor de conductance G_k est :

$$i_k = \frac{G_k}{\sum_{j=1}^N G_j} i \quad (2.15)$$

► **Deux remarques importantes**

- Il faut faire attention à appliquer correctement la formule du diviseur de courant notamment quand on a des associations de résistances en série dans le circuit.
- Il faut également faire attention à l'orientation des intensités.

8. Utilisation de l'équivalence entre les modèles de Thévenin et de Norton

On considère le circuit de la figure 2.25 et on cherche à calculer l'intensité du courant qui traverse la résistance R_c . Sur cet exemple, on va montrer l'efficacité de la méthode utilisant l'équivalence des modèles de Norton et de Thévenin pour simplifier un circuit et en déterminer un modèle équivalent.

On cherche donc un modèle de Thévenin équivalent à la partie gauche du circuit entre les points A et B de la figure 2.25. On note qu'il ne faut surtout pas toucher la résistance R_c .

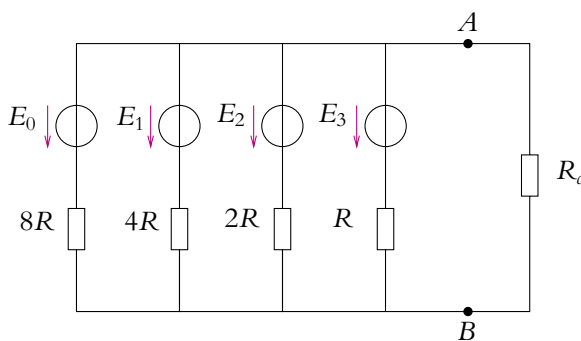


Figure 2.25 Circuit initial.

La règle à suivre, comme pour les méthodes des mailles ou des nœuds, est de transformer les générateurs en modèle de Thévenin (E, r) s'ils sont en série et en modèle de Norton (i_0, g) s'ils sont en parallèles.

Ici tous les générateurs sont en parallèle ; on détermine donc leur modèle de Norton, soit :

Modèle de Thévenin	Modèle de Norton
$(E_0, 8R)$	$\left(\frac{E_0}{8R}, \frac{1}{8R}\right)$
$(E_1, 4R)$	$\left(\frac{E_1}{4R}, \frac{1}{4R}\right)$
$(E_2, 2R)$	$\left(\frac{E_2}{2R}, \frac{1}{2R}\right)$
(E_3, R)	$\left(\frac{E_3}{R}, \frac{1}{R}\right)$

Le circuit obtenu est celui de la figure 2.26.

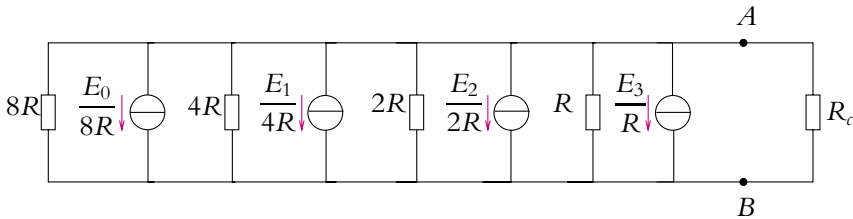


Figure 2.26 Transformation des générateurs en modèles de Norton.

Puisque tous les dipôles sont en parallèle, on ajoute les courants de court-circuit et les conductances pour obtenir le modèle de Norton équivalent à l'ensemble du circuit :

$$\begin{cases} i_{0,\text{tot}} = \frac{E_0}{8R} + \frac{E_1}{4R} + \frac{E_2}{2R} + \frac{E_3}{R} \\ G_{\text{tot}} = \frac{1}{8R} + \frac{1}{4R} + \frac{1}{2R} + \frac{1}{R} = \frac{15}{8R} \end{cases}$$

On obtient le modèle de Thévenin équivalent à l'ensemble du circuit (Cf. figure 2.27) en écrivant :

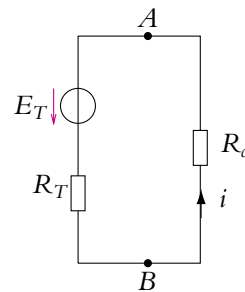


Figure 2.27 Modèle de Thévenin équivalent au circuit.

$$E_T = \frac{i_{0,\text{tot}}}{G_{\text{tot}}} = \frac{1}{15}(E_0 + 2E_1 + 4E_2 + 8E_3) \quad \text{et} \quad R_T = \frac{1}{G_{\text{tot}}} = \frac{8}{15}R \quad (2.16)$$

Il est alors aisé de calculer le courant circulant dans la résistance R_c : il suffit d'appliquer une seule loi des mailles. Ainsi :

$$i = \frac{E_T}{R_T + R_c} = \frac{(E_0 + 2E_1 + 4E_2 + 8E_3)}{8R + 15R_c}$$

Il aurait été bien plus difficile de chercher à calculer ce courant directement à partir du schéma de la figure 2.25 en appliquant les lois de Kirchhoff.

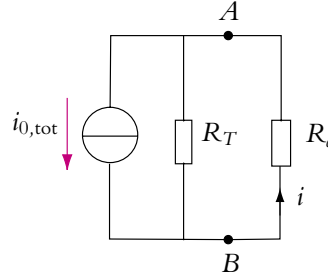


Figure 2.28 Modèle de Norton équivalent au circuit.

► **Remarque :** Pour calculer l'intensité du courant dans la résistance R_c , la dernière étape n'est pas indispensable. En effet, à partir du modèle de Norton du dipôle AB et d'un diviseur de courant, on obtient directement :

$$i = \frac{G_c}{G_c + g_{\text{tot}}} i_{0,\text{tot}} = \frac{\frac{1}{R_c}}{\frac{1}{R_c} + \frac{15}{8R}} \frac{1}{8R} (E_0 + 2E_1 + 4E_2 + 8E_3) = \frac{(E_0 + 2E_1 + 4E_2 + 8E_3)}{8R + 15R_c}$$

On trouve bien évidemment le même résultat.

A. Applications directes du cours

1. Résistance variable et source idéale de tension

On branche en parallèle avec un appareil quelconque une résistance variable R et une source idéale de tension.

1. Comment varie la tension aux bornes de l'appareil quand on fait varier R ?
2. Même question pour l'intensité du courant qui le traverse.
3. Que deviennent ces variations si la source n'est plus idéale ?

2. Associations de dipôles - Utilisation de la caractéristique

Soit un dipôle qu'on suppose caractérisé par une relation $u = f(i)$. La courbe représentative de cette relation est appelée caractéristique.

1. On considère que le dipôle est représenté en convention générateur. Indiquer en justifiant la réponse si le dipôle fonctionne en générateur ou en récepteur suivant le quadrant du plan donnant la tension u en fonction de l'intensité i dans lequel se trouve le point de la caractéristique.
2. Même question si le dipôle est en convention récepteur.
3. Soient deux dipôles D_1 et D_2 , le premier de caractéristique $u = E - Ri$ en convention générateur et le second de caractéristique $u = E' - R'i$ en convention générateur. On relie ces deux dipôles et on note u la tension à leurs bornes et i l'intensité les traversant. Peut-on utiliser pour les deux dipôles une unique convention ? Justifier la réponse.
On prendra dans la suite $E = 3,0 \text{ V}$, $E' = 6,0 \text{ V}$, $R = 1,0 \text{ k}\Omega$ et $R' = 6,0 \text{ k}\Omega$.
4. On suppose dans cette question que D_1 est en convention générateur et D_2 en convention récepteur.
 - a) Déterminer graphiquement puis par le calcul la tension u et l'intensité i .
 - b) Préciser le caractère générateur ou récepteur de chacun des deux dipôles.
5. On inverse les bornes du dipôle D_2 .
 - a) Quelles en sont les conséquences pour l'équation de la caractéristique de ce dipôle ainsi que pour la convention utilisée pour ce dernier ?
 - b) Déterminer graphiquement puis par le calcul la tension u et l'intensité i .
 - c) Préciser le caractère générateur ou récepteur de chacun des deux dipôles.

3. Calcul de la résistance d'un circuit

On considère le circuit de la figure ci-dessous, chaque fil de liaison a une résistance R . Déterminer la résistance totale entre les points A et B puis entre les points C et D .

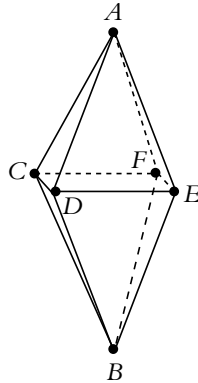


Figure 2.29

4. Résistance équivalente

On considère le montage de la figure ci-dessous. Déterminer l'expression de R' en fonction de R pour que la résistance entre A et B soit égale à R .

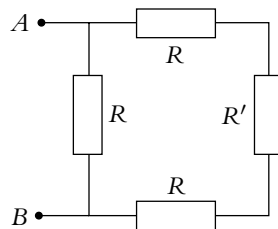


Figure 2.30

5. Détermination d'une intensité par différentes méthodes

On veut déterminer l'intensité parcourant la résistance R_3 du circuit ci-après à partir de plusieurs méthodes :

1. en appliquant des transformations successives entre générateurs de Thévenin et de Norton,
2. en utilisant le théorème de superposition :
 - a) Calculer l'intensité i_1 dans R_3 lorsque seule la source idéale de tension E_1 est allumée.
 - b) En déduire en analysant les symétries et les analogies les intensités i_2 lorsque seule la source idéale de tension E_5 est allumée.
 - c) Retrouver le résultat déjà obtenu par les autres méthodes.

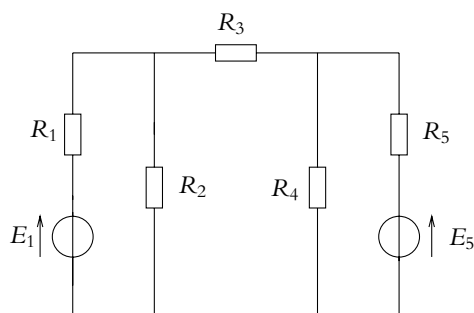


Figure 2.31

6. Résolution par équivalence entre modèles de Thévenin et de Norton

Déterminer l'intensité du courant circulant dans la résistance R du montage suivant en appliquant les équivalences entre modèles de Thévenin et de Norton.

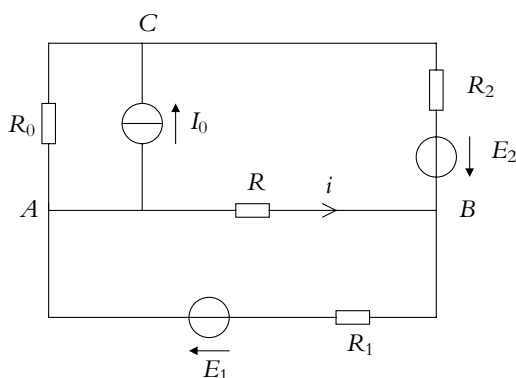


Figure 2.32

7. Circuit en continu, d'après CCP-DEUG 2006

On considère le circuit de la figure ci-dessous. On veut calculer l'intensité i du courant circulant entre A et B .

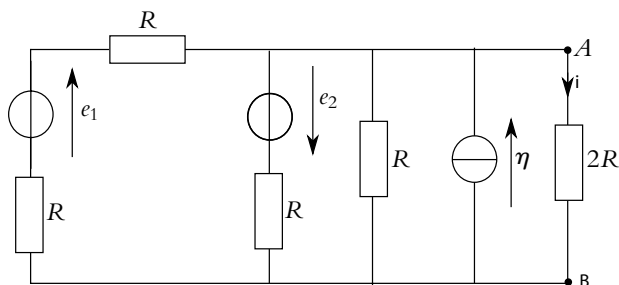


Figure 2.33

1. Effectuer le calcul avec la méthode des noeuds (transformer les modèles de Thévenin en modèles de Norton puis appliquer la loi des noeuds).
2. Effectuer le calcul en déterminant le modèle de Norton du circuit à gauche de AB (circuit moins la résistance $2R$). Utiliser les transformations et associations de générateurs.
3. Effectuer le calcul en déterminant le modèle de Thévenin du circuit à gauche de AB (circuit moins la résistance $2R$). Utiliser les transformations et associations de générateurs.
4. Effectuer le calcul en utilisant le théorème de superposition sans utiliser les transformations de générateur.

8. Autre circuit en continu

On considère le circuit de la figure ci-dessous. On veut calculer l'intensité i du courant circulant entre A et B .

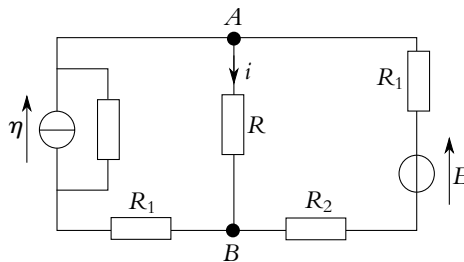


Figure 2.34

1. Déterminer le modèle de Thévenin et le modèle de Norton équivalent au circuit entre A et B en l'absence de R par transformations et associations de générateurs.
2. Déterminer l'intensité du courant dans R .

B. Exercices et problèmes

1. Modélisation simple d'un câble coaxial, d'après CCP-DEUG 2006

Un câble coaxial peut être modélisé par un circuit A_1A_2 , constitué d'une chaîne de n modules identiques comportant chacun trois résistors (résistances respectives $R/2$, $2R$ et $R/2$) (figures 2.35 et 2.36).

Un dipôle résistor X_1X_2 , de résistance $2R$, est branché en parallèle à l'extrémité de la chaîne.

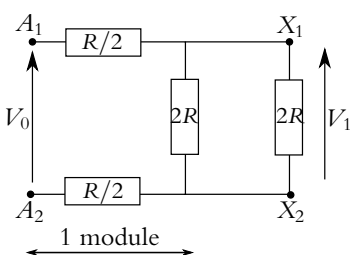


Figure 2.35

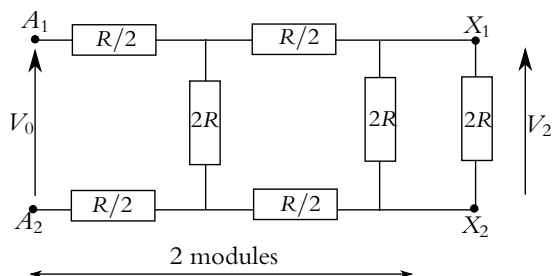


Figure 2.36

1. Le dipôle est équivalent à un résistor.

a) Exprimer, en fonction de R , la résistance équivalente R_1 , dans le cas d'une chaîne ne comportant qu'un seul module (figure 2.35).

b) Même question pour la résistance R_2 , dans le cas d'une chaîne à $n = 2$ modules (figure 2.36).

c) En déduire, sur le même principe, la résistance équivalente R_n d'une chaîne à n modules.

2. Le dipôle A_1A_2 est alimenté par un générateur de f.e.m. constante $V_0 = V_{A1} - V_{A2}$.

a) Déterminer, en fonction de V_0 et R , la tension $V_1 = V_{X1} - V_{X2}$, aux bornes du résistor X_1X_2 , dans le cas d'une chaîne ne comportant qu'un seul module (figure 2.35).

b) Même question pour la tension V_2 , dans le cas d'une chaîne à $n = 2$ modules (figure 2.36).

c) En déduire, sur le même principe, la tension V_n dans le cas d'une chaîne à n modules.

d) En déduire la valeur V_∞ pour une chaîne de longueur infinie ($n \rightarrow \infty$).

2. Transformation triangle-étoile

On considère deux circuits : le circuit triangle (figure 2.37) et le circuit étoile (figure 2.38). On souhaite trouver les relations entre les triplets de résistances (R_1, R_2, R_3) et (R_A, R_B, R_C) pour que, vu de l'extérieur, ces deux circuits aient le même comportement, c'est-à-dire que les tensions V_{AB} , V_{BC} , V_{CA} et les courants i_A , i_B et i_C soient les mêmes dans les deux cas.

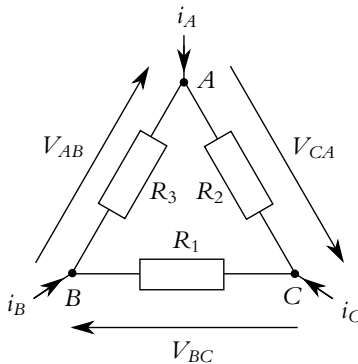


Figure 2.37

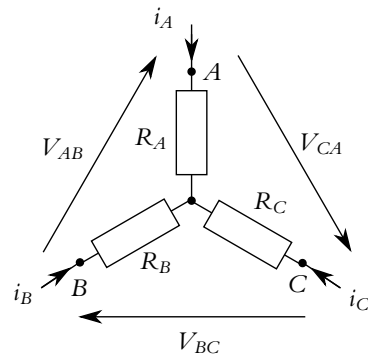


Figure 2.38

On se place dans des cas particuliers, en admettant que, grâce au théorème de superposition, les résultats trouvés seront valables dans tous les cas.

1. Expression de (R_A, R_B, R_C) en fonction de (R_1, R_2, R_3) .

a) On se place dans le cas où $i_A = 0$, $i_B \neq 0$ et $i_C \neq 0$. Représenter les deux circuits équivalents. Par association de résistances, déterminer une relation entre R_B , R_C et (R_1, R_2, R_3) .

- b)** On se place dans le cas où $i_A \neq 0$, $i_B = 0$ et $i_C \neq 0$. Déterminer une relation entre R_A , R_C et (R_1, R_2, R_3) .
- c)** On se place dans le cas où $i_A \neq 0$, $i_B \neq 0$ et $i_C = 0$. Déterminer une relation entre R_A , R_B et (R_1, R_2, R_3) .
- d)** En déduire R_A , R_B , et R_C en fonction de (R_1, R_2, R_3) .

2. Expression de (G_1, G_2, G_3) en fonction de (G_A, G_B, G_C) .

On note G_i la conductance du résistor de résistance R_i .

- a)** On se place dans le cas où $V_{AB} = 0$, $V_{BC} \neq 0$ et $V_{CA} \neq 0$. Représenter les deux circuits équivalents. Par association de conductances déterminer une relation entre G_2 , G_3 et (G_A, G_B, G_C) .
- b)** On se place dans le cas où $V_{AB} \neq 0$, $V_{BC} = 0$ et $V_{CA} \neq 0$. Déterminer une relation entre G_1 , G_3 et (G_A, G_B, G_C) .
- c)** On se place dans le cas où $V_{AB} \neq 0$, $V_{BC} \neq 0$ et $V_{CA} = 0$. Déterminer une relation entre G_1 , G_2 et (G_A, G_B, G_C) .
- d)** En déduire G_1 , G_2 , G_3 en fonction de (G_A, G_B, G_C) .

3. En utilisant la transformation établie précédemment, déterminer la valeur R pour avoir l'équivalence entre les deux montages suivants :

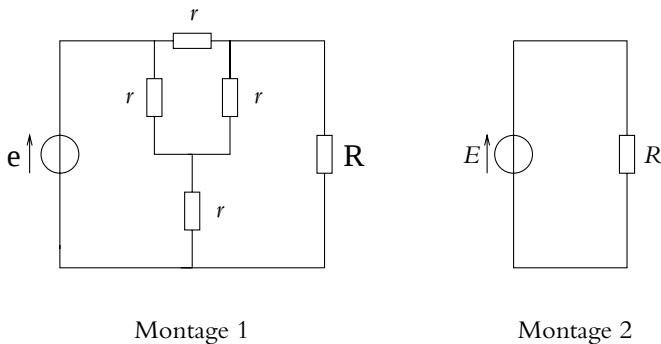


Figure 2.39

3. Comparaison de deux tensions

- 1.** Rappels théoriques :
- a)** Qu'appelle-t-on lois de Kirchhoff?
- b)** Définir une source idéale de courant.
- c)** Rappeler le schéma de principe d'un pont diviseur de tension.
- d)** Etablir la relation du pont diviseur de tension.

2. Méthode d'opposition :

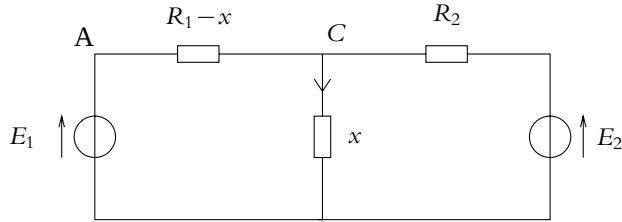


Figure 2.40

- a) On considère le circuit représenté ci-dessus. On notera I_1 l'intensité du courant traversant $R_1 - x$ et I_2 celle du courant traversant R_2 . Peut-on appliquer la relation du pont diviseur de tension pour déterminer la tension aux bornes de x ? Justifier la réponse.
- b) On cherche à calculer toutes les intensités du circuit. Montrer qu'il suffit de résoudre un système de deux équations à deux inconnues pour cela.
- c) Ecrire ce système en prenant I_1 et I_2 comme inconnues.
- d) Le résoudre.
- e) En déduire l'intensité dans x .
- f) On règle la valeur de x pour que I_2 soit nulle. Que vaut alors le rapport $\frac{E_2}{E_1}$?
- g) Justifier qu'on puisse comparer deux tensions à l'aide de ce dispositif.
- h) Que pensez-vous de son utilisation lorsque les tensions sont très différentes?

3. Méthode du diviseur série-parallèle :

On dispose de N résistances R_i pour $i = 1, 2, \dots, N$.

- a) On les associe dans un premier temps en parallèle. Etablir l'expression de la résistance équivalente à cette association.
- b) On place cette association aux bornes d'une source idéale de courant dont on note I_0 le courant de court-circuit. On ajuste la valeur de I_0 pour que la différence de potentiel aux bornes de l'association de résistances soit égale à la tension V_1 et que l'ensemble ne débite pas de courant. Exprimer dans ce cas I_0 en fonction des R_i et de V_1 .
- c) On associe dans un second temps les résistances en série. Etablir l'expression de la résistance équivalente à cette association.
- d) On place l'ensemble aux bornes de la source idéale de courant I_0 dont on ne change pas la valeur. Exprimer la tension V_0 aux bornes des résistances en fonction de V_1 et des résistances R_i .
- e) Exprimer le rapport $\frac{V_0}{V_1}$ en fonction des R_i .
- f) On suppose que toutes les résistances ont la même valeur R . Que vaut alors le rapport $\frac{V_0}{V_1}$?

g) On suppose que V_1 est de l'ordre de 1 V. Comment doit-on choisir le nombre N de résistances si V_0 est de l'ordre de 10^4 V ?

h) En réalité, les résistances ne sont pas rigoureusement égales à R . Pour tenir compte de cet écart, on suppose que $R_i = R + \delta R_i$ avec $\delta R_i \ll R$. On note δR la valeur maximale des écarts. Donner un encadrement à l'ordre le plus bas de $\frac{V_0}{V_1}$ en fonction de δR , R et N . On rappelle que $\frac{1}{1+x} \simeq 1-x$ quand x tend vers 0.

i) En déduire l'erreur maximale commise en ne tenant pas compte des écarts δR_i dans l'expression du rapport des tensions.

j) Expliquer en quoi l'utilisation d'un tel dispositif est intéressante.

4. Convertisseur numérique analogique

Un convertisseur numérique analogique délivre une tension continue U proportionnelle à un nombre décimal N qu'on écrit en base 2.

Dans ce problème, on suppose que $0 \leq N \leq 15$ et que $N = 2^0 a_0 + 2^1 a_1 + 2^2 a_2 + 2^3 a_3$.

On utilise pour cette conversion un réseau $R - 2R$ et quatre interrupteurs K_0 , K_1 , K_2 et K_3 reliés chacun à une source idéale de courant (de courant à vide I_0). On notera $a_i I_0$ le courant délivré avec $a_i = 0$ si l'interrupteur K_i est ouvert et $a_i = 1$ si l'interrupteur K_i est fermé.

Le montage est le suivant :

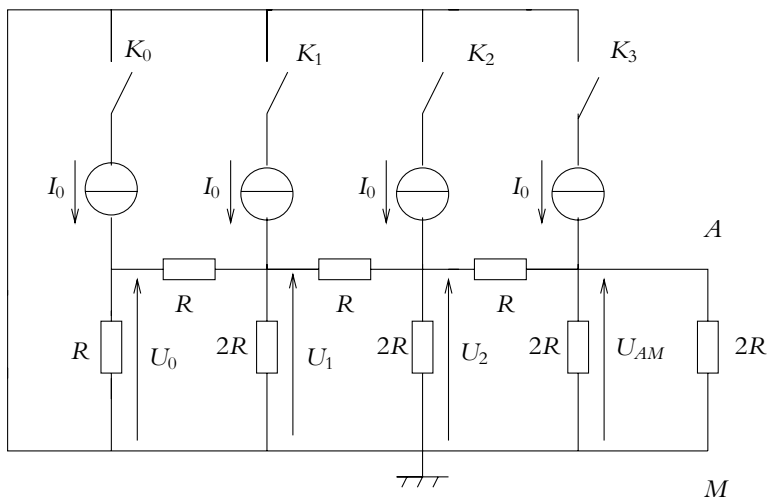


Figure 2.41

1. On suppose dans un premier temps que seul K_3 est fermé. Montrer que le montage est alors équivalent à :

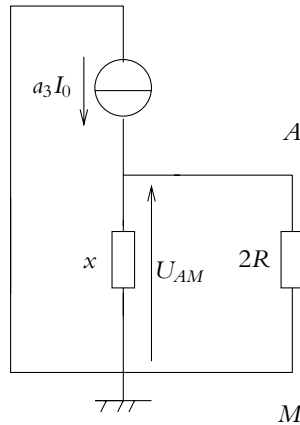


Figure 2.42

On donnera l'expression de x en fonction de R .

2. En déduire l'expression de la tension U_{AM} en fonction de R , a_3 et I_0 .
3. On suppose maintenant que seul K_2 est fermé. Donner le nouveau montage équivalent.
4. Exprimer U_2 en fonction de R , a_2 et I_0 .
5. Déterminer la relation entre U_2 et U_{AM} .
6. En déduire l'expression de la tension U_{AM} en fonction de R , a_2 et I_0 .
7. On suppose maintenant que seul K_1 est fermé. Donner le nouveau montage équivalent.
8. Exprimer U_1 en fonction de R , a_1 et I_0 .
9. Déterminer la relation entre U_1 et U_2 .
10. En déduire celle entre U_1 et U_{AM} .
11. En déduire l'expression de la tension U_{AM} en fonction de R , a_1 et I_0 .
12. On suppose maintenant que seul K_0 est fermé. Donner le nouveau montage équivalent.
13. Exprimer U_0 en fonction de R , a_0 et I_0 .
14. Déterminer la relation entre U_0 et U_1 .
15. En déduire celle entre U_0 et U_{AM} .
16. En déduire l'expression de la tension U_{AM} en fonction de R , a_1 et I_0 .

17. Lorsque tous les interrupteurs peuvent être fermés, déterminer l'expression de U_{AM} en fonction de R , I_0 et des coefficients a_i .

18. Montrer qu'on peut l'écrire sous la forme :

$$U_{AM} = \frac{RI_0}{k} (2^0 a_0 + 2^1 a_1 + 2^2 a_2 + 2^3 a_3)$$

19. Quelle est la valeur minimale de U_{AM} ?

20. Quelle est la plus petite variation possible de U_{AM} ?

21. Quelle est la valeur maximale de U_{AM} ?

Circuits linéaires soumis à un échelon de tension

3

Dans ce chapitre, on étudie les circuits linéaires R , C série, R , L série et R , L , C série soumis à un échelon de tension, c'est-à-dire au régime transitoire de ces circuits entre deux régimes continus. Pour cette étude, on reste dans le cadre de l'approximation des régimes quasi-stationnaires, à savoir qu'on néglige tout phénomène de propagation.

1. Définitions

1.1 Régime continu ou variable

On rappelle que le **régime continu** correspond au fonctionnement d'un circuit lorsque toutes les intensités et toutes les tensions du circuit sont indépendantes du temps : toutes les grandeurs électriques sont des constantes par rapport au temps. Cela sous-entend qu'aucun paramètre des sources n'est modifié.

Par opposition, on parle de **régime variable** quand intensités et tensions varient au cours du temps.

1.2 Régime permanent

On appelle *régime permanent* le fonctionnement où les caractéristiques des intensités et des tensions ne varient pas au cours du temps. Cela englobe le régime continu qui est un régime permanent mais ce n'est pas le seul type de régime permanent. On verra en deuxième période que les signaux sinusoïdaux correspondent également à des régimes permanents. Il sera donc nécessaire de bien distinguer un régime permanent d'un régime continu et de ne pas faire d'amalgame entre ces deux notions.

1.3 Régime transitoire

Cependant il est rare qu'un phénomène dure toute l'éternité : il y a un instant où par exemple on allume les sources. On peut donc passer d'un régime permanent à un autre.

On appelle *régime transitoire* le régime durant lequel on passe d'un régime permanent à un autre.

C'est par exemple ce qui se passe à l'établissement ou à l'arrêt d'une source de courant ou de tension. Il s'agit d'un régime temporaire par opposition aux précédents régimes. Malgré son caractère éphémère, il est cependant d'une grande importance. Il peut en effet se produire des phénomènes très brefs durant lesquels intensité et/ou tension risquent de prendre des valeurs très grandes, ce qui risque d'endommager les composants en les soumettant à des contraintes en dehors de leur domaine de fonctionnement.

On s'intéresse donc dans ce chapitre à ce type de régime sur le cas simple des circuits électriques linéaires ne comprenant que des résistors, des bobines, des condensateurs et des générateurs modélisables à l'aide des modèles de Thévenin ou de Norton.

1.4 Échelon de tension ou de courant

On s'intéresse en particulier au passage brusque d'un régime continu à un autre régime continu. Cela correspond à la réponse du système qu'on appelle un échelon de tension ou de courant.

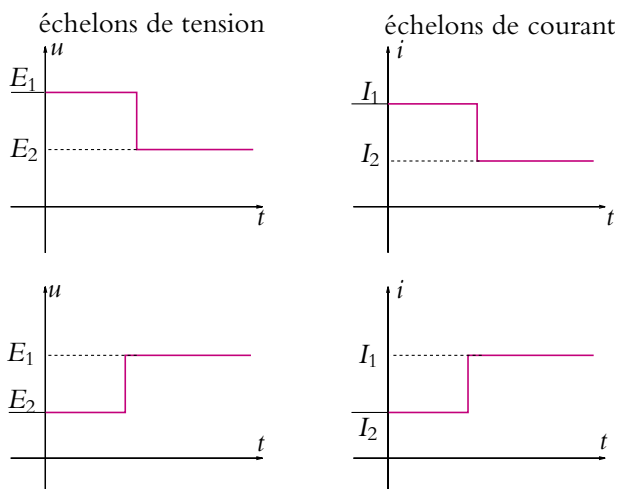


Figure 3.1 Échelons de tension et de courant.

Un échelon de tension (respectivement de courant) est le passage brusque d'une tension continue (respectivement d'un courant continu) à une autre tension continue

(respectivement à un autre courant continu). Le laps de temps pour passer de l'un à l'autre est infinitésimal.

On se limite dans la suite aux échelons de tension.

2. Circuits du premier ordre

2.1 Équation différentielle du circuit R, C

On considère le circuit constitué d'un résistor de résistance R en série avec un condensateur de capacité C , l'ensemble étant soumis à une tension $e(t)$:

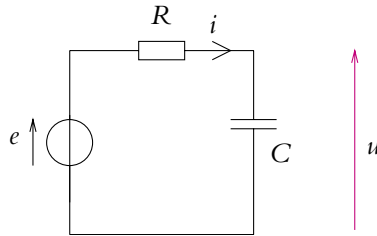


Figure 3.2 Circuit R, C .

On cherche à déterminer l'expression de l'intensité $i(t)$ et de la tension $u(t)$.

La loi des mailles donne :

$$e = Ri + u$$

Or la relation entre l'intensité traversant un condensateur et la tension à ses bornes s'écrit :

$$i = C \frac{du}{dt}$$

On obtient alors l'équation différentielle suivante :

$$\frac{du}{dt} + \frac{1}{RC}u = \frac{1}{RC}e$$

soit

$$\frac{du}{dt} + \frac{u}{\tau} = \frac{e}{\tau} \quad (3.1)$$

en posant $\tau = RC$.

Il s'agit d'une équation différentielle du premier ordre à coefficients constants dont la solution s'écrit comme la somme :

- de la solution générale de l'équation homogène associée (c'est-à-dire à second membre nul),
- d'une solution particulière.

On verra dans ce chapitre les solutions particulières correspondant au cas où le second membre est constant au cours du temps. En deuxième période, on s'intéressera aux solutions particulières dans le cas des régimes sinusoïdaux forcés. Dans tous les cas, les solutions particulières décrivent le régime permanent.

On s'intéresse dans un premier temps à la solution générale de l'équation homogène associée qui traduit le régime libre.

2.2 Régime libre

On appelle *régime libre* la réponse d'un système en l'absence d'excitation. En pratique, l'expression de cette réponse est la solution de l'équation différentielle homogène associée à savoir de l'équation différentielle à second membre nul.

En reprenant l'exemple donné au paragraphe précédent, cette solution s'écrit :

$$u = U_0 \exp\left(-\frac{t}{\tau}\right)$$

où U_0 est une constante à déterminer à partir des conditions initiales.

2.3 Charge et décharge d'un condensateur

On s'intéresse maintenant à la réponse à un échelon de tension.

a) Charge d'un condensateur

On suppose tout d'abord que $e(t)$ passe d'une valeur 0 pour $t < 0$ à une valeur E pour $t > 0$.

On doit tenir compte ici à la fois :

- du régime libre décrit par la solution générale de l'équation homogène associée,
- du régime permanent donné par $u = E$ dont on vérifie qu'elle est bien une solution particulière de l'équation différentielle.

La constante U_0 qui a été introduite dans l'expression de la solution correspondant au régime libre doit être déterminée à partir de la solution générale de l'équation différentielle (3.1) c'est-à-dire de la somme de la solution générale de l'équation homogène associée et d'une solution particulière décrivant le régime permanent.

On utilise alors la propriété importante que **la tension d'un condensateur est une fonction continue du temps**. On a justifié cette propriété lors de l'étude énergétique d'une capacité par le fait que l'énergie stockée dans un condensateur est proportionnelle au carré de sa charge ou de la tension à ses bornes et qu'une discontinuité de l'énergie correspondrait à une puissance infinie, ce qui n'a pas de réalité physique. Alors en supposant que le condensateur est initialement déchargé, on a :

$u(0) = U_0 + E = 0$ soit $U_0 = -E$ et finalement la solution s'écrit :

$$u = E \left(1 - \exp \left(-\frac{t}{\tau} \right) \right)$$

L'intensité peut s'en déduire :

$$i = C \frac{du}{dt} = CE \times \frac{1}{\tau} \exp \left(-\frac{t}{\tau} \right) = \frac{E}{R} \exp \left(-\frac{t}{\tau} \right)$$

Les allures de la tension et de l'intensité en fonction du temps sont donc :

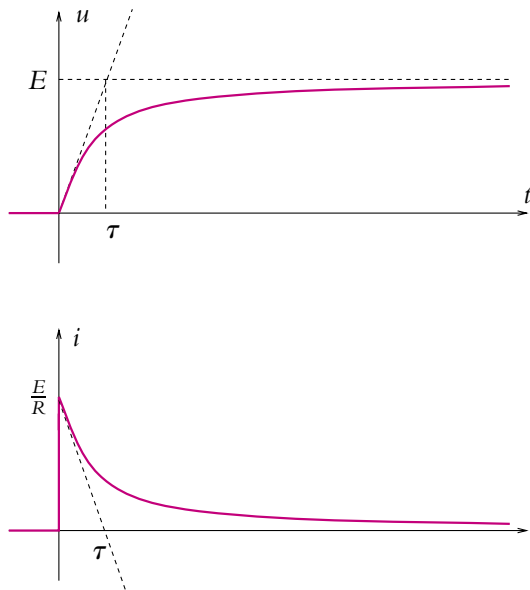


Figure 3.3 Charge d'un condensateur.

On peut noter que la tension aux bornes du condensateur tend vers une valeur limite E de façon exponentielle.

D'autre part, l'intensité présente une discontinuité en $t = 0$: elle passe de la valeur 0 à la valeur $\frac{E}{R}$. Physiquement, cette discontinuité peut s'expliquer par la structure même des condensateurs : le courant arrive sur l'isolant et ne le traverse pas instantanément.

b) Décharge d'un condensateur

Il s'agit de l'effet inverse : on suppose que le condensateur possède une charge initiale $Q = CE$. Ceci est rendu possible par l'existence d'une valeur limite (obtenue au paragraphe précédent) de la tension aux bornes du condensateur. À $t = 0$, on ouvre le circuit précédent, ce qui revient à faire passer la tension e de la valeur E à la valeur 0.

L'équation différentielle est simplement $\frac{du}{dt} + \frac{u}{\tau} = 0$ dont la solution est :

$$u = U_1 \exp\left(-\frac{t}{\tau}\right)$$

avec, par continuité de u , $u(0) = U_1 = E$ soit

$$u = E \exp\left(-\frac{t}{\tau}\right)$$

On en déduit l'intensité :

$$i = C \frac{du}{dt} = -\frac{E}{R} \exp\left(-\frac{t}{\tau}\right)$$

Les allures de l'intensité et de la tension au cours du temps sont donc les suivantes :

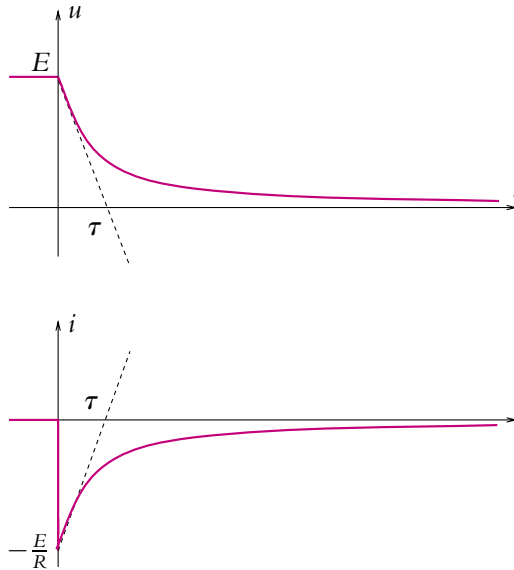


Figure 3.4 Décharge d'un condensateur.

On remarque que la tension aux bornes du condensateur tend vers 0, ce qui justifie le nom de décharge du condensateur.

c) Temps caractéristique

Dans les deux cas de charge ou de décharge du condensateur, on voit apparaître une constante de temps :

$$\tau = RC$$

En effet, le produit RC est homogène à un temps puisque l'exponentielle ne peut agir que sur un nombre sans dimension. On va chercher à interpréter cette constante de temps ou temps caractéristique.

Pour cela, on détermine la tangente à la courbe à l'origine. On se limite au cas de la charge et on laisse le soin au lecteur d'établir des résultats analogues pour la décharge.

Pour la tension, l'équation de la tangente à l'origine s'écrit :

$$u = \left(\frac{du}{dt} \right)_{t=0} t = \frac{E}{\tau} t$$

Pour l'intensité, on a :

$$i = \left(\frac{di}{dt} \right)_{t=0} t + i(t=0) = \frac{E}{R} \left(1 - \frac{t}{\tau} \right)$$

On cherche l'intersection de la tangente à l'origine avec l'asymptote.

- pour la tension, on doit résoudre :

$$\frac{E}{\tau} t = E$$

dont la solution est $t = \tau$;

- pour l'intensité, on résout :

$$\frac{E}{R} \left(1 - \frac{t}{\tau} \right) = 0$$

dont la solution est aussi $t = \tau$.

La durée τ s'interprète donc comme le temps qui serait nécessaire pour atteindre la valeur finale limite avec une croissance linéaire au lieu d'être exponentielle à partir des mêmes conditions initiales. Cette propriété reste vraie à partir d'un instant t_0 quelconque ($t_0 \geq 0$)

On en déduit que plus le temps caractéristique est grand, plus la tangente à l'origine a une pente faible et plus le système est lent à atteindre le régime permanent.

Au bout de quelques τ , on peut considérer :

- que la tension u est égale à E et que l'intensité i est nulle lors de la charge du condensateur,
- que la tension u et l'intensité i sont nulles lors de la décharge du condensateur.

d) Aspect énergétique

Pour ce paragraphe, on se place dans le cas où la charge et la décharge sont complètes.

L'étude énergétique d'un tel circuit peut se faire à partir de la loi des mailles :

$$R \frac{dq}{dt} + \frac{q}{C} = e.$$

En multipliant cette relation par $idt = \frac{dq}{dt} dt$, on obtient :

$$Ri^2 dt + \frac{1}{C} q \frac{dq}{dt} dt = ei$$

soit

$$Ri^2 dt + d\left(\frac{q^2}{2C}\right) = eidt$$

L'interprétation des différents termes donne :

- $eidt$ est l'énergie fournie par la source de tension,
- $Ri^2 dt$ correspond à l'énergie dissipée par effet Joule dans le résistor R lors du passage du courant entre d et dt ,
- $d\left(\frac{q^2}{2C}\right)$ correspond à l'énergie reçue ou « stockée » par le condensateur sous forme d'énergie électrostatique.

Pendant la charge, le générateur fournit l'énergie

$$W = \int_0^{\infty} eidt = E \int_0^Q dq = EQ = CE^2$$

Le condensateur stocke l'énergie :

$$W_C = \int_0^{+\infty} \frac{d}{dt} \left(\frac{q^2}{2C} \right) dt = \int_0^Q d \left(\frac{q^2}{2C} \right) = \frac{Q^2}{2C} = \frac{CE^2}{2}$$

Le condensateur stocke donc la moitié de l'énergie fournie par le générateur et l'autre moitié est dissipé par effet Joule dans le résistor.

Pendant la décharge, il n'y a plus d'énergie fournie : le générateur est éteint.

Le condensateur « stocke » l'énergie :

$$W_C = \int_0^{+\infty} \frac{d}{dt} \left(\frac{q^2}{2C} \right) dt = \int_Q^0 d \left(\frac{q^2}{2C} \right) = -\frac{Q^2}{2C} = -\frac{CE^2}{2}$$

Le condensateur restitue l'énergie qu'il a stocké lors de la charge. Cette énergie est dissipé par effet Joule dans le résistor, en l'absence de toute autre utilisation.

$$W_J = \int_0^{+\infty} Ri^2 dt = -W_C$$

► **Remarque :** On retiendra qu'un condensateur peut stocker et restituer de l'énergie alors qu'un résistor ne fait que dissiper l'énergie sous forme de chaleur par effet Joule.

e) Réponse à un signal carré

Les résultats précédents peuvent s'appliquer au cas où l'entrée est un signal carré de période T .

On commence par faire l'hypothèse qu'à $t = 0$, la source de tension passe de 0 à E et que la tension aux bornes du condensateur est nulle. On observe donc la charge du condensateur étudiée au paragraphe a). À $t = \frac{T}{2}$, le signal d'entrée revient à 0 et on se trouve dans la situation de décharge du condensateur.

La différence peut provenir de la définition de la condition initiale. En effet, le condensateur risque de ne pas être totalement chargé au moment du basculement. Il faudra alors prendre la tension aux bornes du condensateur à cet instant à savoir $u\left(\frac{T}{2}\right) = E\left(1 - \exp\left(-\frac{T}{2\tau}\right)\right)$ et non la valeur E . On obtient alors $u = E\left(\exp\left(+\frac{T}{2\tau}\right) - 1\right)\exp\left(-\frac{t}{\tau}\right)$.

À $t = T$, on aura un nouveau basculement de 0 à E et une nouvelle charge du condensateur. Cependant celle-ci risque de ne pas être totalement déchargée et il faudra prendre

$$E(T) = E\left(\exp\left(-\frac{T}{2\tau}\right) - \exp\left(-\frac{T}{\tau}\right)\right)$$

comme nouvelle condition initiale au lieu de 0.

Et ainsi de suite. On obtient donc le chronogramme de la figure 3.5.

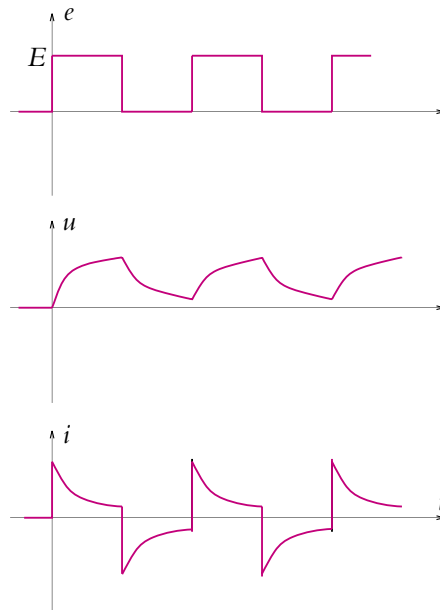


Figure 3.5 Réponse à un créneau lorsque le régime permanent peut s'établir.

Pour les deux cas limites $T \ll \tau$ et $T \gg \tau$, les chronogrammes sont les suivants :

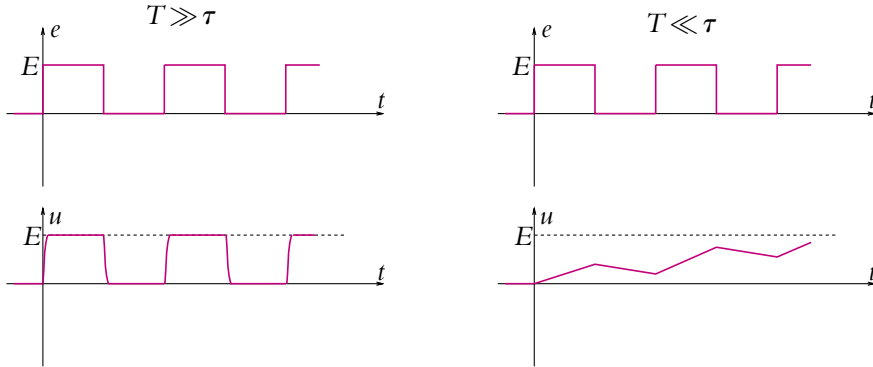


Figure 3.6 Influence de la période sur la réponse à un créneau.

2.4 Circuit R, L - Établissement et rupture du courant dans une bobine

a) Équation différentielle

On considère le circuit constitué d'un résistor de résistance R en série avec une bobine d'inductance L , l'ensemble étant soumis à une tension $e(t)$ qui passe d'une valeur nulle à une valeur E à l'instant $t = 0$:

La loi des mailles donne :

$$e = Ri + L \frac{di}{dt}$$

et on obtient l'équation différentielle suivante :

$$\frac{di}{dt} + \frac{R}{L}i = \frac{1}{L}e$$

Il s'agit d'une équation différentielle du premier ordre à coefficients constants du même type que celle obtenue lors de l'étude du circuit R, C . On peut l'écrire sous la forme :

$$\frac{di}{dt} + \frac{i}{\tau} = \frac{e}{R\tau}$$

en notant $\tau = \frac{L}{R}$ qu'on interprète comme un temps caractéristique.

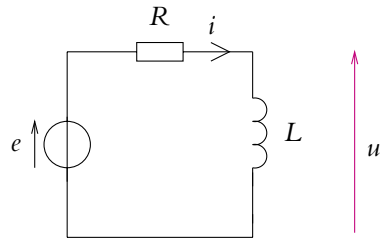


Figure 3.7 Circuit R, L .

La solution pour $t > 0$ s'écrit comme la somme :

- de la solution générale de l'équation homogène associée :

$$i = I_0 \exp\left(-\frac{t}{\tau}\right)$$

- d'une solution particulière :

$$i = \frac{E}{R}$$

car $e(t) = E$ est constant pour tout instant $t > 0$.

b) Établissement du courant dans une bobine

Il s'agit du cas analogue à la charge du condensateur, on en déduit donc l'expression de l'intensité :

$$i = \frac{E}{R} \left(1 - \exp\left(-\frac{t}{\tau}\right)\right)$$

Pour déterminer la constante I_0 , on utilise cette propriété importante : **l'intensité qui traverse une bobine est une fonction continue du temps**. Cette propriété a déjà été justifiée par la continuité de l'énergie stockée par la bobine.

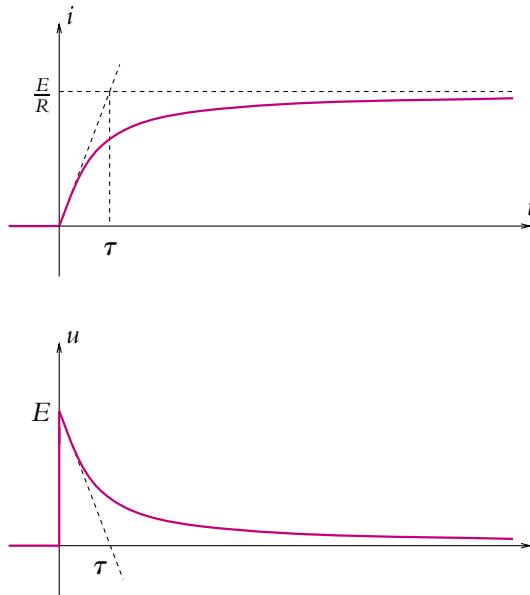


Figure 3.8 Établissement du courant dans une bobine.

L'expression de la tension s'obtient par dérivation :

$$u = L \frac{di}{dt} = E \exp\left(-\frac{t}{\tau}\right)$$

Les allures de la tension et de l'intensité sont données ci-dessus.

On peut noter que l'intensité qui circule dans la bobine tend vers une valeur limite $\frac{E}{R}$ de façon exponentielle.

D'autre part, la tension passe en $t = 0$ de la valeur 0 à la valeur E : elle présente donc une discontinuité aux bornes de la bobine.

c) Rupture du courant dans une bobine

Il s'agit de l'effet inverse : on suppose que la tension $e(t)$ passe de la valeur E à 0. On résout l'équation différentielle et on en déduit l'expression suivante de l'intensité :

$$i = \frac{E}{R} \exp\left(-\frac{t}{\tau}\right)$$

et de la tension :

$$u = L \frac{di}{dt} = L \frac{E}{R} \left(-\frac{R}{L} \exp\left(-\frac{t}{\tau}\right) \right)$$

Là encore, on utilise le fait que l'intensité du courant qui traverse une bobine est une fonction continue du temps.

Les allures de l'intensité et de la tension au cours du temps sont donc les suivantes :

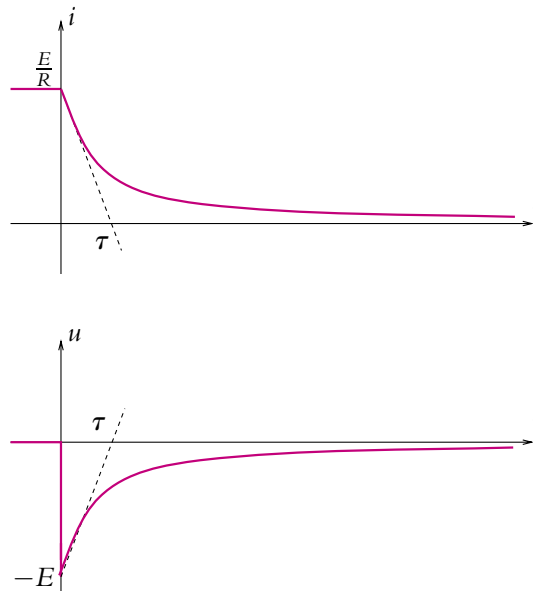


Figure 3.9 Rupture du courant dans une bobine.

On remarque que l'intensité dans la bobine tend vers 0, ce qui justifie le nom de rupture du courant dans la bobine.

d) Aspect énergétique

Pour ce paragraphe, on se place dans le cas où l'établissement et la rupture sont complets.

L'étude énergétique d'un tel circuit peut se faire directement à partir de l'équation différentielle écrite pour l'intensité i :

$$L \frac{di}{dt} + Ri = e$$

En multipliant cette relation par $idt = \frac{dq}{dt} dt$, on obtient :

$$d \left(\frac{Li^2}{2} \right) + Ri^2 dt = e idt$$

soit

$$Li \frac{di}{dt} dt + Ri^2 dt = e idt$$

L'interprétation des différents termes donne :

- $e idt$ est l'énergie fournie par la source de tension entre t et $t + dt$,
- Ri^2 correspond à l'énergie dissipée par effet Joule dans le résistor R lors du passage du courant entre t et $t + dt$,
- $d \left(\frac{Li^2}{2} \right)$ correspond à l'énergie reçue ou « stockée » par dans la bobine sous forme d'énergie magnétique.

Pendant l'établissement du courant, le générateur fournit l'énergie

$$W = \int_0^\infty e idt = \frac{LE^2}{R^2} = LI_0^2$$

La bobine stocke une partie de l'énergie :

$$W_L = \int_0^{+\infty} \frac{d}{dt} \left(\frac{Li^2}{2} \right) dt = \int_0^{I_0} d \left(\frac{Li^2}{2} \right) = \frac{LE^2}{2R^2} = \frac{LI_0^2}{2}$$

Le reste est dissipé par effet Joule dans le résistor.

Pendant la rupture du courant, il n'y a plus d'énergie fournie : le générateur est éteint.

La bobine « stocke » l'énergie :

$$W_L = \int_0^{+\infty} \frac{d}{dt} \left(\frac{Li^2}{2} \right) dt = \int_{I_0}^0 d \left(\frac{Li^2}{2} \right) = -\frac{LE^2}{2R^2} = -\frac{LI_0^2}{2}$$

La bobine restitue l'énergie qu'elle a stockée lors de l'établissement du courant. Cette énergie est alors dissipée par effet Joule dans le résistor, en l'absence de toute autre utilisation.

On retiendra qu'une bobine peut stocker et restituer de l'énergie alors qu'un résistor ne fait que dissiper l'énergie sous forme de chaleur par effet Joule.

e) Réponse à un signal carré

Les résultats des paragraphes 2.3.b et 2.3.c peuvent s'appliquer au cas où l'entrée est un signal carré de période T .

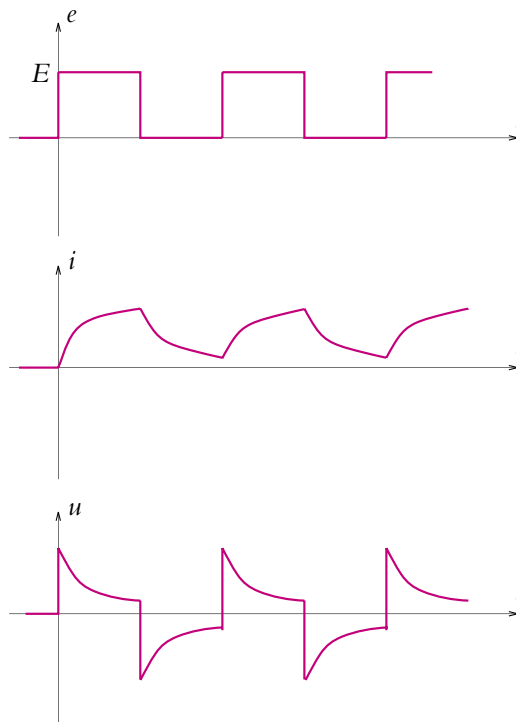


Figure 3.10 Réponse d'un circuit R, L à un crêteau.

Par analogie à l'étude faite lors de la charge et de la décharge d'un condensateur à travers un résistor, on obtient le chronogramme de la figure ci-dessus pour $T \simeq \tau$.

3. Circuits du second ordre

3.1 Équation différentielle du circuit R, L, C série (figure 3.11)

On applique une loi des mailles :

$$u + Ri + L \frac{di}{dt} = e$$

Or la relation tension - intensité pour un condensateur s'écrit :

$$i = C \frac{du}{dt}$$

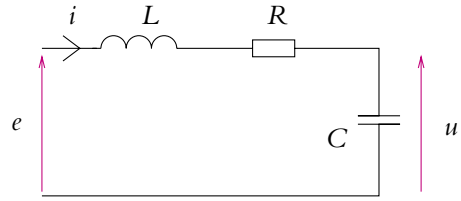


Figure 3.11 Circuit R, L, C série.

On obtient finalement l'équation différentielle suivante :

$$LC \frac{d^2 u}{dt^2} + RC \frac{du}{dt} + u = e$$

Il s'agit d'une équation différentielle linéaire du second ordre à coefficients constants et second membre non nul.

3.2 Régime libre ou propre

Comme précédemment, on appelle régime propre ou régime libre le régime correspondant à l'absence de source extérieure.

On cherche donc les solutions de l'équation différentielle homogène associée :

$$LC \frac{d^2 u}{dt^2} + RC \frac{du}{dt} + u = 0$$

qui est une équation différentielle du second ordre à coefficients constants sans second membre.

On peut remarquer qu'on peut réécrire cette équation différentielle en utilisant les paramètres $\omega_0 = \frac{1}{\sqrt{LC}}$ et $\xi = \frac{R}{2} \sqrt{\frac{C}{L}}$. On verra que ω_0 est homogène à une pulsation et s'exprime donc en rad.s^{-1} et que ξ est une grandeur sans dimension.

$$\frac{d^2 u}{dt^2} + \frac{R}{L} \frac{du}{dt} + \frac{1}{LC} u = 0$$

soit

$$\frac{d^2 u}{dt^2} + \frac{R\sqrt{C}\omega_0}{\sqrt{L}} \frac{du}{dt} + \omega_0^2 u = 0$$

et

$$\frac{d^2 u}{dt^2} + 2\xi\omega_0 \frac{du}{dt} + \omega_0^2 u = 0$$

a) Équation caractéristique

Ce type d'équations différentielles se résout en introduisant la notion d'équation caractéristique. Il s'agit de chercher des solutions sous la forme $u = U_0 \exp(rt)$ donc :

$$(r^2 + 2\xi\omega_0 r + \omega_0^2) U_0 \exp(rt) = 0$$

Comme $\exp(rt)$ ne s'annule jamais et que $U_0 = 0$ conduit à la solution identiquement nulle ne présentant pas d'intérêt, on en déduit l'équation du second degré en r appelée équation caractéristique :

$$r^2 + 2\xi\omega_0 r + \omega_0^2 = 0$$

Le discriminant s'écrit :

$$\Delta = 4\omega_0^2 (\xi^2 - 1)$$

qui peut prendre tous les signes possibles. On obtient les trois cas suivants :

1. $\Delta > 0$: l'équation caractéristique admet deux solutions réelles :

$$r = -\xi\omega_0 \pm \omega_0 \sqrt{\xi^2 - 1}$$

2. $\Delta = 0$: l'équation caractéristique admet une solution double :

$$r = -\xi\omega_0$$

3. $\Delta < 0$: l'équation caractéristique admet deux solutions complexes conjuguées :

$$r = -\xi\omega_0 \pm j\omega_0 \sqrt{1 - \xi^2}$$

où $j^2 = -1$.

On peut s'interroger sur la signification physique de la valeur limite. On aura $\Delta > 0$ à condition que $C^2 \left(R^2 - 4\frac{L}{C} \right) > 0$ soit $R > 2\sqrt{\frac{L}{C}}$. La dissipation d'énergie ayant lieu au niveau de la résistance, cette inégalité permet de mesurer l'importance de la dissipation d'énergie dans le système.

Les trois cas conduisent à trois types de solutions différentes.

b) Régime apériodique

Il correspond au cas où $\Delta > 0$ soit $R > 2\sqrt{\frac{L}{C}}$ ou $\xi > 1$. On a alors une forte dissipation par effet Joule du fait d'une résistance importante.

L'équation caractéristique admet deux solutions réelles :

$$r_1 = -\xi\omega_0 - \omega_0 \sqrt{\xi^2 - 1} \quad \text{et} \quad r_2 = -\xi\omega_0 + \omega_0 \sqrt{\xi^2 - 1}$$

La solution de l'équation différentielle est une combinaison linéaire des deux fonctions $\exp(r_1 t)$ et $\exp(r_2 t)$:

$$u = U_1 \exp(r_1 t) + U_2 \exp(r_2 t)$$

où U_1 et U_2 sont deux constantes à déterminer à l'aide des conditions initiales. En effet, $u(t)$ est la tension aux bornes d'un condensateur, c'est donc une fonction continue du temps, tout comme $i(t) = C \frac{du}{dt}$ qui est l'intensité du courant traversant une bobine.

On peut tout d'abord noter que r_1 et r_2 sont négatifs. Pour r_1 , le résultat est évident à cause des signes « - » et du fait que ξ et ω_0 sont par définition des quantités positives. Pour r_2 , il faut remarquer que $\xi^2 > \xi^2 - 1$, ce qui implique la négativité de r_2 . Ce résultat assure l'existence physique de la solution qui ne tend jamais vers l'infini. Une solution tendant vers l'infini n'aurait pas de sens physique et un tel résultat sous-entendrait que d'autres phénomènes non linéaires doivent être pris en compte.

De la solution donnant la tension u aux bornes du condensateur, on peut déduire l'expression de l'intensité i (ce qui permettrait, si le besoin s'en faisait sentir, d'obtenir les tensions aux bornes de l'inductance et de la résistance) :

$$i(t) = C \frac{du}{dt} = Cr_1 U_1 \exp(r_1 t) + Cr_2 U_2 \exp(r_2 t)$$

En supposant que U_1 et U_2 sont positifs et avec comme conditions initiales $u(0) = U$ et $i(0) = 0$, on obtient les allures données par la figure ci-contre pour la tension u et l'intensité i . On notera qu'ici on utilise à la fois la continuité de l'intensité dans la bobine et la continuité de la tension aux bornes du condensateur pour calculer les valeurs de U_1 et U_2 .

Il ne s'agit que d'une allure. Il faudrait réaliser une étude plus complète de la fonction pour garantir le tracé. Cependant on peut noter les points suivants :

- la solution tend vers 0 à l'infini,
- les conditions initiales imposent la valeur initiale ainsi que la pente initiale des courbes,
- il n'y a pas d'oscillations mais uniquement des combinaisons linéaires d'exponentielles décroissantes : les seuls paramètres modifiant l'allure ne peuvent être que les constantes U_1 et U_2 .

L'absence d'oscillations donne le nom d'*apériodique* à ce régime caractérisé par l'absence de phénomènes périodiques.

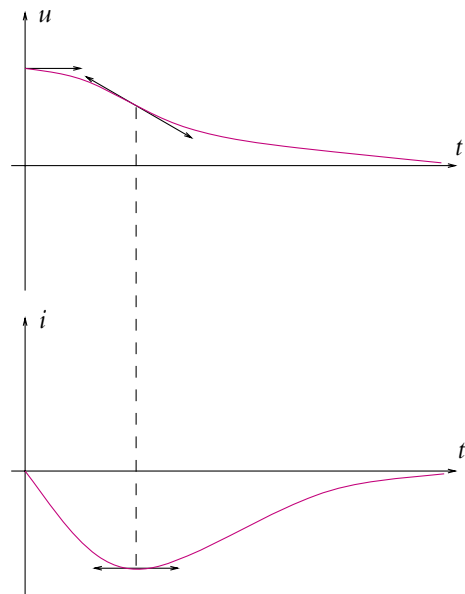


Figure 3.12 Régime apériodique.

c) Régime critique

Le *régime critique* correspond à la valeur $\Delta = 0$ soit $R = 2\sqrt{\frac{L}{C}}$ et $\xi = 1$ qui est la valeur critique entre le régime apériodique qui vient d'être étudié et le régime pseudopériodique qui sera vu au paragraphe suivant.

Dans ce cas, l'équation caractéristique admet une solution double

$$r = -\omega_0$$

La solution peut donc s'écrire sous la forme :

$$u(t) = (U_0 + U_1 t) \exp(-\omega_0 t)$$

où U_0 et U_1 sont des constantes à déterminer avec les conditions initiales (grâce à la continuité de u et de i à $t = 0$).

On en déduit l'intensité dans la branche :

$$i = C \frac{du}{dt} = C (U_1 - \omega_0 U_0 - \omega_0 U_1 t) \exp(-\omega_0 t)$$

L'allure de $u(t)$ et $i(t)$ sont du même type que pour le régime apériodique. Les mêmes remarques relatives aux allures peuvent être formulées.

Il n'y a donc pas de différence majeure au niveau du sens physique des régimes apériodique et critique. On peut noter qu'expérimentalement le régime critique sera impossible à obtenir car il faudra que la condition $R = 2\sqrt{\frac{L}{C}}$ soit rigoureusement vérifiée, ce qui est en pratique impossible à réaliser. La seule différence provient du fait que les fonctions tendent vers 0 plus vite dans le cas du régime critique. Le retour à l'équilibre est alors plus rapide pour le régime critique que pour le régime apériodique.

Il faut retenir que le régime critique est un cas limite entre les deux autres régimes.

d) Régime pseudopériodique

Il correspond au cas où le discriminant Δ de l'équation caractéristique est négatif soit $R < 2\sqrt{\frac{L}{C}}$ ou $\xi < 1$, ce qui correspond à une faible dissipation et à une faible valeur de la résistance. On a donc deux solutions complexes conjuguées :

$$r_{\pm} = -\xi \omega_0 \pm j \omega_0 \sqrt{1 - \xi^2} = -\xi \omega_0 \pm j \omega$$

La solution de l'équation différentielle peut s'écrire sous la forme :

$$u(t) = (U_1 \cos \omega t + U_2 \sin \omega t) \exp(-\xi \omega_0 t) = U_0 \cos(\omega t + \varphi) \exp(-\xi \omega_0 t)$$

où U_1 et U_2 (ou bien U_0 et φ) sont des constantes que l'on détermine à l'aide des conditions initiales, en utilisant la continuité des fonction $u(t)$ et $i(t)$.

On obtient donc deux parties dans cette solution :

- des oscillations périodiques à la pulsation

$$\omega = \omega_0 \sqrt{1 - \xi^2}$$

correspondant à la partie imaginaire des solutions de l'équation caractéristique,

- une atténuation exponentielle de l'amplitude associée à la partie réelle des solutions de l'équation caractéristique.

C'est pourquoi on qualifie ce régime de *pseudopériodique* ou encore de sinusoïdal amorti.

Dans la suite, on utilisera la solution sous la forme : $u(t) = U_0 \cos(\omega t + \varphi) \exp(-\xi \omega_0 t)$.

On peut en déduire l'expression de l'intensité par la relation :

$$\begin{aligned} i(t) &= C \frac{du}{dt} \\ &= -CU_0 (\alpha \omega_0 \cos(\omega t + \varphi) + \omega \sin(\omega t + \varphi)) \exp(-\xi \omega_0 t) \end{aligned}$$

On en déduit l'allure suivante sur laquelle on peut formuler les remarques suivantes :

- elles tendent toujours vers 0 en oscillant quelles que soient les conditions initiales,
- les conditions initiales imposent les valeurs à l'origine des courbes ainsi que leur tangente,
- un maximum pour la tension u correspond à une valeur nulle pour l'intensité puisque $i = C \frac{du}{dt}$,
- les courbes sont comprises entre deux exponentielles qui en constituent les enveloppes.

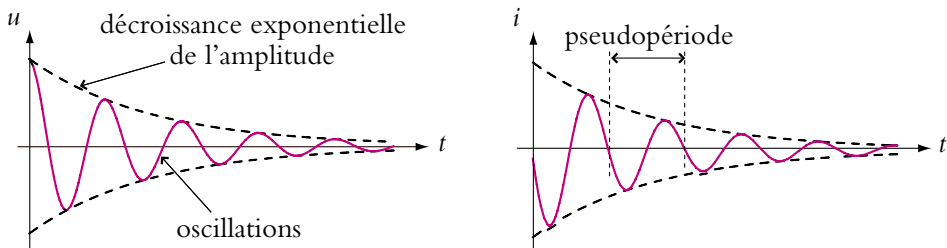


Figure 3.13 Régime pseudopériodique.

Pseudopériode

On définit une *pseudopériode* correspondant à la période de la partie oscillante en $\cos \omega t$. Il s'agit donc de

$$T = \frac{2\pi}{\omega} = \frac{2\pi}{\omega_0 \sqrt{1 - \xi^2}}$$

Si on avait un régime sinusoïdal non amorti (correspondant à l'absence de terme en $\frac{d\mu}{dt}$ et donc à une résistance R nulle), on aurait une période dite période propre :

$$T_0 = \frac{2\pi}{\omega_0}$$

Décrément logarithmique

On définit une grandeur permettant de mesurer la rapidité de la décroissance exponentielle de l'amplitude : il s'agit du *décrément logarithmique*. Pour cela, on considère deux instants successifs pour lesquels l'amplitude passe par un maximum, c'est-à-dire pour lesquels $\cos(\omega t + \varphi) = 1$. Il s'agit de deux instants séparés par une pseudopériode.

Les amplitudes correspondantes sont respectivement

$$u(t + nT) = U \exp(-\xi \omega_0 (t + nT))$$

et

$$u(t + (n + 1)T) = U \exp(-\xi \omega_0 (t + (n + 1)T)).$$

Le rapport entre ces deux amplitudes vaut :

$$\begin{aligned} \frac{u(t + nT)}{u(t + (n + 1)T)} &= \frac{U \exp(-\xi \omega_0 (t + nT))}{U \exp(-\xi \omega_0 (t + (n + 1)T))} \\ &= \exp(\xi \omega_0 T) < 1 \end{aligned}$$

La décroissance de la tension d'une pseudopériode à l'autre correspond à une suite géométrique qu'on quantifie comme suit :

$$\begin{aligned} \delta &= \ln \frac{u(t + nT)}{u(t + (n + 1)T)} \\ &= \xi \omega_0 T \\ &= \frac{2\pi\xi}{\sqrt{1 - \xi^2}} \end{aligned}$$

δ est appelé décrément logarithmique.

Cas d'un faible amortissement

Il s'agit d'un régime pseudopériodique pour lequel l'exponentielle $\exp(-\xi \omega_0 t)$ vaut presque 1 car $\xi \omega_0 \ll 1$. On aura alors une très faible valeur de ξ et de la résistance. Dans ce cas, la pseudopériode s'identifie à la période propre :

$$T \simeq T_0$$

e) Comparaison des trois régimes

Les résultats précédents peuvent être résumés par la figure suivante :

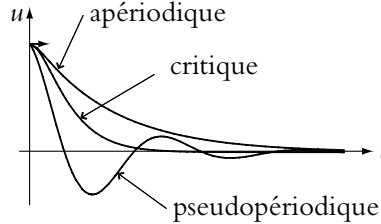


Figure 3.14 Les trois régimes d'un circuit R, L, C série.

3.3 Réponse à un échelon

Dans ce cas, l'équation différentielle a un second membre non nul, ce dernier étant constant. La solution est donc la somme :

- de la solution de l'équation homogène associée ou régime libre $u_{\text{libre}}(t)$ qui vient d'être déterminée,
- d'une solution particulière de l'équation avec second membre, que l'on cherche constante.

On aura alors si $e = E_0$ une solution particulière $u = E_0$ et la solution globale s'écrit :

$$u(t) = E_0 + u_{\text{libre}}(t)$$

On retrouve pour $u_{\text{libre}}(t)$ l'un des trois cas précédents.

Le point important à noter est la nécessaire détermination de la valeur des constantes de la solution $u_{\text{libre}}(t)$ en tenant compte de la solution particulière. Les conditions initiales sont définies globalement et s'appliquent à la solution générale même si les constantes n'apparaissent que dans la partie $u_{\text{libre}}(t)$. La procédure est donc la suivante :

1. déterminer la forme de la solution de l'équation homogène associée en introduisant des constantes,
2. chercher une solution particulière,
3. déterminer la valeur des constantes à partir des conditions initiales sur la solution générale.

3.4 Aspects énergétiques

Comme pour les circuits du premier ordre, on reprend l'équation différentielle pour réaliser l'étude énergétique. On l'écrit :

$$e = L \frac{di}{dt} + Ri + \frac{q}{C}$$

et on la multiplie par $idt = \frac{dq}{dt}dt$ soit :

$$\begin{aligned} eidt &= Li \frac{di}{dt} dt + Ri^2 dt + \frac{q}{C} \frac{dq}{dt} dt \\ &= Ri^2 dt + \left(\frac{Li^2}{2} + \frac{q^2}{2C} \right) \end{aligned}$$

par des transformations déjà vues dans la partie sur les circuits du premier ordre.

Ces différents termes peuvent s'interpréter de la façon suivante.

1. $eidt$ est l'énergie fournie par le générateur entre t et $t + dt$,
2. Ri^2 est l'énergie dissipée par effet Joule dans le résistor sous forme de chaleur entre t et $t + dt$,
3. $d \left(\frac{Li^2}{2} + \frac{q^2}{2C} \right)$ correspond à la variation d'énergie entre t et $t + dt$:
 - magnétique $\frac{Li^2}{2}$ de la bobine,
 - électrostatique $\frac{q^2}{2C}$ du condensateur.

Cette énergie est stockée ou cédée lors du passage d'un courant dans le circuit.

Le bilan énergétique montre que l'énergie apportée par le générateur est en partie dissipée par effet Joule dans le résistor et en partie stockée dans la bobine et dans le condensateur. Ces derniers peuvent à leur tour céder cette énergie quand il n'y a plus de générateur, cette énergie cédée est alors dissipée dans la résistance.

3.5 Facteur de qualité dans le cas d'un amortissement très faible

Dans le cas d'un amortissement très faible, à savoir d'une très faible valeur de résistance, on peut calculer la perte d'énergie au cours d'une pseudopériode T .

L'énergie $\mathcal{E}$ est une fonction quadratique de l'intensité ou de la tension. Chacune de ces deux grandeurs est, dans le cas d'un amortissement très faible, une grandeur oscillante dont l'amplitude diminue exponentiellement avec un facteur $\exp(-\xi\omega_0 t)$. L'énergie oscille donc également et le facteur de décroissance vaut : $\exp(-2\xi\omega_0 t)$. Pour deux instants séparés par une pseudopériode, la phase sera la même. On en déduit le rapport :

$$\frac{\mathcal{E}(t+T)}{\mathcal{E}(t)} = \exp(-2\xi\omega_0 T)$$

et :

$$\frac{\Delta\mathcal{E}}{\mathcal{E}} = \frac{\mathcal{E}(t) - \mathcal{E}(t+T)}{\mathcal{E}(t)} = 1 - \frac{\mathcal{E}(t+T)}{\mathcal{E}(t)} = 1 - \exp(-2\xi\omega_0 T)$$

Lorsque l'amortissement est faible, on a $\xi \rightarrow 0$ et $T \simeq T_0$ soit $\xi\omega_0 T \simeq \xi\omega_0 T_0$ et on peut faire un développement limité :

$$\frac{\Delta\mathcal{E}}{\mathcal{E}} = 1 - (1 - 2\xi\omega_0 T_0 + \dots) \simeq 2\xi\omega_0 T_0 = 4\pi\alpha$$

On pose $Q = \frac{2}{\xi}$. Il s'agit du facteur de qualité du circuit oscillant qu'on reverra en seconde période. On obtient donc ici une interprétation énergétique du facteur de qualité dans le cas d'un amortissement très faible et uniquement dans ce cas :

$$Q = 2\pi \frac{\mathcal{E}}{\Delta\mathcal{E}}$$

A. Applications directes du cours

1. Annulation d'un régime transitoire

On considère un circuit formé de l'association en série d'un générateur de tension idéale de f.e.m. E , d'une bobine réelle d'inductance L et de r , et d'un interrupteur.

1. Déterminer l'expression du courant parcourant le circuit après la fermeture de l'interrupteur. Interpréter les deux termes obtenus.
2. On souhaite annuler dans le courant traversant le générateur, les effets du régime transitoire lié à la fermeture de l'interrupteur. Pour cela, on place en parallèle sur la bobine un condensateur de capacité C en série avec une résistance R .
 - a) Déterminer la nouvelle expression du courant traversant le générateur.
 - b) En déduire les valeurs de C et R à choisir.

2. Influence d'un condensateur dans un circuit RL

Soit un générateur de tension de force électromotrice E et de résistance interne r . Il est placé aux bornes d'une bobine d'inductance L et de résistance R . On suppose qu'à l'instant initial, l'ensemble fonctionne en régime permanent. On branche alors en parallèle avec la bobine un condensateur de capacité C . L'objet de ce problème est l'étude de l'intensité i traversant la bobine.

1. Déterminer la valeur de i à $t = 0^+$ juste après avoir branché le condensateur.
2. Même question pour $\frac{di}{dt}$.
3. Établir l'équation différentielle vérifiée par i .
4. Rappeler la valeur numérique usuelle de r .
5. On donne $L = 43 \text{ mH}$, $R = 9,1 \, \Omega$ et $E = 5,0 \text{ V}$. Quelle valeur doit-on prendre pour C pour observer un régime quasipériodique ?
6. Déterminer l'expression littérale de i si $C = 10 \, \mu\text{F}$.

3. Deux condensateurs en série avec une résistance, d'après I.C.N.A. 2003

Deux condensateurs de capacités respectives C_1 et C_2 sont reliés par une résistance R (figure 3.15). À l'instant initial, leurs charges respectives sont $Q_{10} = Q_0$ et $Q_{20} = 0$. On pose : $\tau = R \frac{C_1 C_2}{C_1 + C_2}$

1. Établir l'expression à la date t de l'intensité i du courant dans la résistance R .
2. Déterminer les charges $Q_1(t)$ et $Q_2(t)$ des deux condensateurs à la date t .

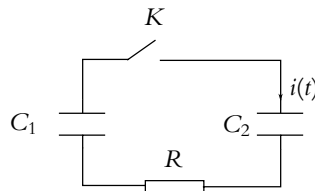


Figure 3.15

3. Calculer la variation de l'énergie emmagasinée dans les deux condensateurs entre l'instant initial t_i et l'instant final t_f .
4. Calculer l'énergie consommée par effet Joule dans la résistance R entre les instants $t = 0$ et t .

4. Condensateur en parallèle sur une bobine et alimenté par un échelon de courant

Une source de courant idéale délivre un échelon de courant :

- $i_g(t) = 0$ pour $t \leq 0$.
- $i_g(t) = I_0$ pour $t > 0$.

dans un circuit constitué d'un condensateur de capacité C , initialement déchargé, en parallèle sur une bobine d'inductance L et de résistance R , avec $R < 2\sqrt{\frac{L}{C}}$.

1. Établir l'équation différentielle du courant $i(t)$ dans la bobine :

$$\frac{d^2 i}{dt^2} + \frac{R}{L} \frac{di}{dt} + \frac{1}{LC} i = \frac{1}{LC} I_0$$

Mettre cette équation sous forme canonique et déterminer l'expression de la pulsation propre ω_0 et du facteur de qualité Q .

2. Déterminer l'expression générale de $i(t)$ faisant apparaître deux constantes.
3. a) Quelles sont les grandeurs continues à $t = 0$?
b) En déduire les valeurs de $i(0^+)$ et $u(0^+)$ (tension aux bornes de la bobine L, R), à $t = 0^+$.
c) Déterminer les constantes d'intégration de $i(t)$.

B. Exercices et problèmes

1. Circuit RL alimenté par une tension créneau, d'après ENSTIM 2003

On se propose d'étudier la réponse d'un circuit (RL) à une tension en créneaux délivrée par un générateur basse fréquence (G.B.F.). Le circuit représenté sur la figure ci-dessous comporte une bobine parfaite d'inductance L , un résistor de résistance R et un G.B.F. délivrant une tension en créneaux u représentée sur la figure ci-dessous.

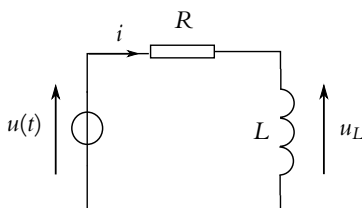


Figure 3.16

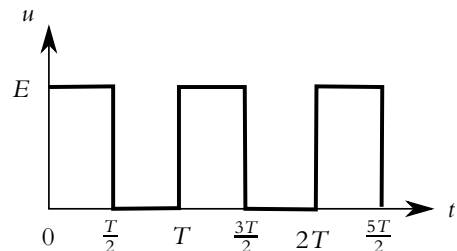


Figure 3.17

- On définit la constante de temps τ , exprimée en secondes, du circuit (RL) par une relation du type $\tau = L^\alpha R^\beta$ où α et β sont deux constantes réelles. Par analyse dimensionnelle rapide, déterminer la valeur des exposants α et β (On raisonnera à partir des relations entre u et i).
- Pour $0 \leq t \leq \frac{T}{2}$, établir l'équation différentielle régissant les variations de l'intensité i dans le circuit. La résoudre en justifiant soigneusement la détermination de la (des) constante(s) d'intégration.
 - En déduire l'expression de $u_L(t)$.
 - Tracer l'allure des courbes représentatives de $i(t)$ et de $u_L(t)$ en précisant les valeurs vers lesquelles ces fonctions tendent en régime permanent, ainsi que l'équation des tangentes à l'origine.
- Déterminer complètement l'expression de $i(t)$ et $u_L(t)$ pour $\frac{T}{2} \leq t \leq T$.
- Le G.B.F. est réglé sur la fréquence $f = 1,0$ kHz, la bobine a pour inductance $L = 1,0$ H et $R = 1,0 \cdot 10^3 \Omega$. Comparer la période T de la tension délivrée par le G.B.F. et la constante de temps τ du circuit. Tracer qualitativement l'évolution des graphes de $i(t)$ et $u_L(t)$ sur quelques périodes. Commenter.

2. Régime transitoire d'un circuit R, L, C parallèle, d'après Mines d'Alès

Soit un circuit constitué de l'association en série d'une source idéale de tension E , d'une résistance R , d'un interrupteur K et de l'association en parallèle d'une résistance r , d'une inductance L et d'une capacité C . Initialement l'interrupteur est ouvert, la capacité est déchargée et tous les courants sont nuls. On ferme l'interrupteur K à l'instant $t = 0$. On notera i l'intensité du courant dans R , i_1 dans L , i_2 dans C et i_3 dans r ainsi que u la tension aux bornes de r (ou de C ou de L).

- Déterminer, en les justifiant, les valeurs de u , i , i_1 , i_2 et i_3 juste après la fermeture de l'interrupteur K .
- Même question lorsque le régime permanent est complètement établi.
- Établir l'équation différentielle vérifiée par i_3 pour $t > 0$.
- L'écrire sous la forme

$$\frac{d^2 i_3}{dt^2} + 2\lambda\omega_0 \frac{di_3}{dt} + \omega_0^2 i_3 = 0$$

On donnera les expressions de ω_0 et λ en fonction de r , R , L et C .

- Établir la condition que doivent vérifier r , R , L et C pour observer un régime pseudo-périodique.
- Que caractérise λ ?
- Définir la pseudo-pulsation ω et la pseudo-période T et en donner les expressions en fonction de r , R , L et C .
- Déterminer l'expression de i_3 en fonction du temps. On justifiera la détermination des constantes.
- Définir et exprimer le décrétement logarithmique en fonction de T , λ et ω_0 .

10. Déterminer le temps nécessaire pour que le régime permanent soit établi dans le circuit avec une précision d'un millièm.
11. Déterminer, en les justifiant, les valeurs de u , i , i_1 , i_2 et i_3 juste après l'ouverture de l'interrupteur K .
12. Établir l'équation différentielle vérifiée par i_3 lorsqu'on ouvre l'interrupteur. On ne se préoccupera pas de ce qui se passe juste à l'instant de l'ouverture.
13. L'écrire sous une forme analogue à celle de la question 4.
14. Quelle est la nature du régime ? On suppose que la condition de la question 5. est vérifiée.
15. A quelle condition aura-t-on un signal non nul lors de l'ouverture de l'interrupteur ?

3. Circuit RC à deux mailles, d'après ICNA 2006

On considère le montage de la figure ci-dessous. On ferme l'interrupteur à $t = 0$. Avant la fermeture, le condensateur C_a est déchargé alors que la tension aux bornes de C_b est V_0 .

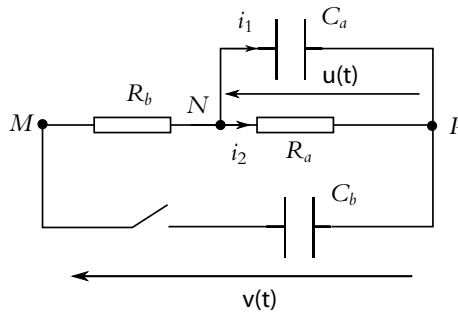


Figure 3.18

1. Établir les relations :

$$C_b \frac{dV}{dt} = \frac{u - V}{R_b} \quad \text{et} \quad i = C_a \frac{du}{dt} + \frac{u}{R_a}$$

2. Montrer que la tension u vérifie l'équation différentielle :

$$\frac{d^2 u}{dt^2} + \frac{\omega_0}{Q} \frac{du}{dt} + \omega_0^2 u = 0$$

où Q et ω_0 sont des constantes à déterminer en fonction de C_a , C_b , R_a et R_b .

3. a) Montrer que les deux solutions de l'équation caractéristique sont des réels négatifs que l'on notera r_1 et r_2 .
- b) En déduire l'expression générale de $u(t)$.
- c) Montrer qu'à $t = 0^+$, $i_2 = 0$ et en déduire que $i_1 = V_0/R_b$.
- d) En déduire l'expression finale de $u(t)$.

4. Étude d'un flash, d'après Centrale TSI 2006

Attention, l'ouverture d'un appareil photo jetable est dangereuse car le circuit électrique comporte un condensateur pouvant être chargé jusqu'à une tension de 300 V.

Le schéma électrique du circuit commandant le flash est fortement simplifié. Il est constitué d'un transformateur dont on va étudier le circuit primaire puis le circuit secondaire.

L'alimentation est une pile de f.e.m. continue $E = 1,5 \text{ V}$. Pour atteindre des tensions de 300 V il faut utiliser un transformateur, lequel ne fonctionne qu'avec des tensions alternatives.

La première partie du montage comporte donc un circuit RLC qui permet d'obtenir une tension variable à partir d'une tension continue.

Le circuit étudié, dont l'interrupteur est fermé lorsqu'on arme le flash (à l'instant $t = 0$) est représenté par la figure (3.19). La résistance R_i est la résistance interne de la pile. Le condensateur est initialement déchargé.

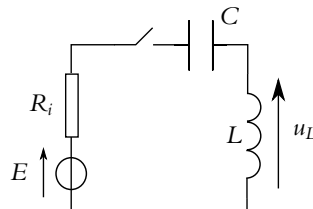


Figure 3.19

- Déterminer la valeur de u_L au bout d'un temps très long.
- Que vaut u_L à l'instant $t = 0^+$, c'est-à-dire juste après la fermeture de l'interrupteur ?
- Déterminer l'équation différentielle vérifiée par u_L au cours du temps.
- Résoudre cette équation en négligeant la résistance R_i (on prendra $R_i = 0$).
- Pour cette question, on ne néglige plus la résistance R_i . Évaluer le temps nécessaire pour que u_L atteigne la valeur déterminée à la question (1). On prendra $R_i = 0,5 \Omega$, $C = 200 \text{ pF}$ et $L = 36 \text{ mH}$.

Le schéma équivalent au circuit secondaire est représenté sur la figure (3.20). La tension au secondaire du transformateur est alternative et son amplitude est 200 fois plus grande que celle au primaire. Après transformation en tension continue U_0 (redressement), cette tension permet de charger un condensateur de capacité C' jusqu'à la tension U_0 . Lors de la prise de la photo, l'interrupteur se ferme et le condensateur se décharge dans le flash qui est alors équivalent à une résistance R_f .

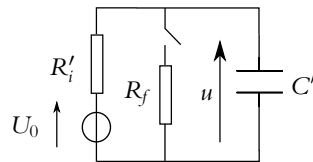


Figure 3.20

- Déterminer l'énergie contenue initialement (avant fermeture de l'interrupteur) dans le condensateur. Effectuer l'application numérique avec $C' = 150 \mu\text{F}$ et $U_0 = 300 \text{ V}$.
- Déterminer le modèle de Thévenin (E_t , R_t) équivalent à U_0 , R_i' et R_f . Application numérique pour $R_f = 10 \Omega$ et $R_i' = 180 \Omega$.
- Déterminer l'évolution de u au cours du temps.
- Déterminer l'ordre de grandeur du temps nécessaire à la décharge du condensateur.

5. Étude de circuits R, L , d'après Banque Agro 2004

1. Circuit simple :

On considère dans un premier temps le circuit représenté ci-dessous. A l'instant $t = 0$, on ferme l'interrupteur K qui était ouvert depuis très longtemps.

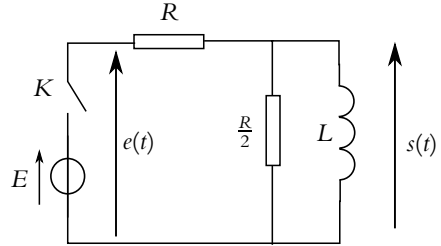


Figure 3.21

a) L'intensité du courant i traversant la résistance R est-elle continue en $t = 0$? Si non, donner ses valeurs en 0^- et en 0^+ .

b) Mêmes questions pour la tension s .

c) Que vaut $s(t)$ lorsque t tend vers l'infini?

d) Établir l'équation différentielle vérifiée par $s(t)$.

e) En déduire l'expression de $s(t)$.

f) Tracer l'allure de $s(t)$.

g) Exprimer en fonction de L et de R le temps t_0 au bout duquel la tension s a été divisée par 10.

h) Proposer une méthode expérimentale pour déterminer t_0 à l'aide d'un oscilloscope. On précisera le montage à utiliser et la méthode de la mesure pratique.

i) On mesure $t_0 = 3,0 \mu\text{s}$ pour $R = 1000 \Omega$. En déduire la valeur de L .

j) On remplace le générateur idéal de tension par un générateur délivrant un créneau de période T . Quel est l'ordre de grandeur de la fréquence à utiliser pour pouvoir effectivement mesurer t_0 par la méthode proposée ci-dessus?

2. Circuit plus complexe :

On considère maintenant le circuit ci-dessous. L'interrupteur K est ouvert depuis très longtemps et on le ferme à l'instant $t = 0$.

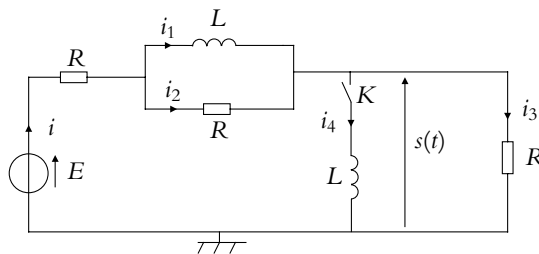


Figure 3.22

a) Déterminer les valeurs des intensités i_4 , i_1 , i , i_2 et i_3 à l'instant $t = 0^+$.

b) Déterminer les valeurs de la tension s et des intensités i_2 , i , i_1 , i_3 et i_4 lorsque t tend vers l'infini.

c) Déterminer les relations entre s et i_3 puis entre s et i_4 .

d) Établir la relation entre R , L , s , $\frac{ds}{dt}$ et $\frac{di}{dt}$.

- e) Même question pour R , L , i_2 , $\frac{di_2}{dt}$ et $\frac{di}{dt}$.
- f) Déterminer la relation entre $\frac{dE}{dt}$, R , L , s , i_2 et $\frac{ds}{dt}$.
- g) Établir la relation entre R , L , s , $\frac{ds}{dt}$, $\frac{d^2s}{dt^2}$, $\frac{dE}{dt}$ et $\frac{d^2E}{dt^2}$.
- h) En déduire l'équation différentielle vérifiée par s sachant que le générateur de tension est idéal de tension à vide E .
- i) En déduire l'expression de $s(t)$ en ne cherchant pas à déterminer les constantes d'intégration.
- j) Donner une relation entre les deux constantes d'intégration.

6. Circuit alimenté par un hacheur

Un hacheur est un système électronique qui permet d'obtenir une tension continue réglable à partir d'une alimentation constante E . Dans le cas présent le hacheur est modélisé par un interrupteur K (figure 3.23) qui peut basculer entre la position (1) et la position (2). Le condensateur est supposé de capacité C suffisamment grande pour que la tension V_c puisse être considérée constante.

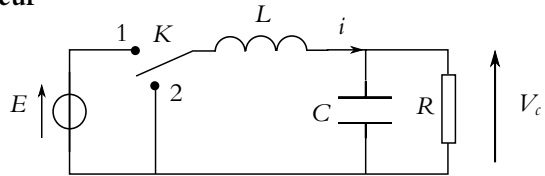


Figure 3.23

L'interrupteur K a les positions suivantes :

- K est en (1) pour $0 < t < \alpha T$;
- K est en (2) pour $\alpha T < t < T$.

α est un nombre positif inférieur à 1 et T la période du hacheur.

- a) Dans un premier temps, K est position (1). Établir l'équation différentielle reliant l'intensité i (figure 3.23) à E , V_c et L .

b) On note I_m l'intensité à $t = 0$. Déterminer l'expression de $i(t)$ en fonction des données.

c) Déterminer l'expression de $i(\alpha T)$, on note I_M cette valeur.
- a) À l'instant $t = \alpha T$, l'interrupteur bascule en position (2). Établir la nouvelle équation différentielle à laquelle obéit i .

b) Déterminer l'expression de $i(t)$ pour $\alpha T \leq t \leq T$ en fonction de α , T , E , V_c , et I_m .

c) L'intensité i est périodique en régime établi, et donc $i(T) = i(0) = I_m$. Calculer V_c en fonction de E et α .
- On rappelle que la valeur moyenne d'une grandeur $s(t)$ est $\langle s \rangle = \frac{1}{T} \int_0^T s(t) dt$.

a) Montrer que la valeur moyenne de l'intensité dans C est nulle. Calculer la valeur moyenne de i et en déduire l'expression de $I_M + I_m$ en fonction de E , α et R .

b) Déterminer I_m puis I_M en fonction de α , E , R , L et T .

4. Quelle valeur maximale peut prendre R pour que I_m soit positif?

Mécanique 1

Cinématique du point matériel

4

1. Définitions

1.1 Notion de point matériel

On se limite en première période à l'étude mécanique des *points matériels*; en deuxième période, on s'intéressera également aux systèmes de deux points matériels.

Il convient donc de définir la notion de point matériel et d'analyser rapidement les conditions pouvant permettre d'identifier les systèmes réels à un ou plusieurs points matériels.

a) Définition d'un solide

Avant de préciser la définition d'un point matériel, il est nécessaire de donner celle d'un solide.

On appelle *solide* tout système matériel pour lequel les distances entre deux points du système sont constantes et invariantes au cours du temps. On peut noter qu'il s'agit d'un modèle simplificateur : un solide subit des déformations liées aux contraintes qui lui sont imposées. Par exemple, lorsqu'on tire sur les deux extrémités d'une barre métallique, celle-ci peut se déformer, voire se casser, si la contrainte est trop importante. La définition qui vient d'être donnée exclut ce type de déformations : cela suppose que celles-ci sont négligeables par rapport aux autres aspects mécaniques.

De manière générale, la position d'un solide sera déterminée à l'aide de six paramètres :

- les trois coordonnées d'un point du solide (on choisira souvent le centre d'inertie du système),
- trois paramètres angulaires définissant l'orientation d'un trièdre ($OXYZ$) lié au solide (au sens où ($OXYZ$) est immobile par rapport à ce dernier) par rapport au trièdre ($Oxyz$) définissant le référentiel dans lequel on travaille.

b) Définition d'un point matériel

On appelle *point matériel* un solide dont la position est entièrement définie par la seule donnée des trois coordonnées d'un point du solide. Cela revient à négliger tout effet de rotation du solide sur lui-même ou son extension spatiale.

c) Validité du concept de point matériel

Si on considère par exemple un ballon, peut-on le considérer comme un point matériel ?

Dans le cas d'un ballon de rugby, sa rotation sur lui-même est le plus souvent visible du fait de sa forme ovoïdale. Il paraît alors difficile d'assimiler le ballon de rugby à un point ; le seul cas envisageable est celui d'une translation au cours de laquelle le ballon ne tourne pas sur lui-même.

Dans le cas d'un ballon de football, sa forme sphérique ne permet pas de visualiser les effets liés à la rotation du solide sur lui-même. On peut alors étudier sa trajectoire comme celle d'un point matériel. La rotation du ballon, que l'on peut considérer comme uniforme au cours du temps, intervient pourtant dans l'expression de l'énergie cinétique du solide et influence donc son mouvement.

Cependant, dans un certain nombre de situations, l'approximation d'un solide par un point matériel permet d'interpréter correctement la trajectoire observée. Il s'agit du cas où la rotation du solide sur lui-même est négligeable.

Par ailleurs, même lorsque cette approximation n'est pas possible, le concept de point matériel sera utile. En effet, l'étude du mouvement du centre d'inertie d'un solide s'effectue dans ce cadre. De même, le mouvement d'un système quelconque (comme un solide ou un fluide) pourra être analysé en décomposant par la pensée le système en petits éléments matériels qu'on pourra assimiler à des points matériels du fait notamment de leur faible extension spatiale.

1.2 Espace et distance

Il est nécessaire pour repérer la position de points matériels de doter l'espace d'une mesure et donc d'une unité. Celle-ci est le mètre dont la définition légale a évolué au cours du temps.

- Avant 1960, on utilisait un mètre-étalon constitué par la distance entre deux traits gravés sur une barre de platine iridié (toujours conservée au Bureau International des Poids et Mesures au Pavillon de Breteuil à Sèvres).
- Entre 1960 et 1983, il s'agissait d'un multiple de la longueur d'onde d'une radiation émise par l'atome de krypton lorsqu'un électron passe d'un niveau d'énergie à un autre.

- La définition actuelle est la longueur du trajet parcouru dans le vide par la lumière pendant une durée de $\frac{1}{299792458}$ seconde. En effet, la vitesse de la lumière est fixée conventionnellement à $299792458 \text{ m.s}^{-1}$: il s'agit d'une constante universelle.

1.3 Temps et durée

Le temps est le paramètre permettant de dater un événement et d'établir la chronologie d'une succession d'événements. Ceci est rendu possible par l'hypothèse d'uniformité du temps : les lois physiques sont invariantes par translation dans le temps.

Il faut alors choisir une unité de temps : il s'agira de la seconde. Comme celle du mètre, la définition de la seconde a évolué au cours du temps.

- Avant 1960, la seconde correspondait à $\frac{1}{86400}$ de la durée de rotation de la Terre sur elle-même appelée « jour sidéral ».
- De 1960 à 1967, on a utilisé le temps des éphémérides lié à la rotation de la Terre autour du Soleil.
- Actuellement la seconde est légalement définie comme la durée de $9\,192\,631\,770$ périodes de la radiation correspondant à la transition entre deux raies hyperfines de l'état fondamental de l'isotope¹ 133 du césium. L'interaction entre l'électron et le champ magnétique créé par le noyau conduit à distinguer pour un même niveau d'énergie (par exemple le fondamental) plusieurs sous-niveaux qui constituent la structure hyperfine de l'atome. Lorsqu'un électron passe d'un niveau hyperfin à un autre, une radiation est émise, caractérisée par sa période. La seconde est donc un multiple de la période de la transition hyperfine de l'état fondamental de l'atome de césium 133. Il s'agit du temps atomique.

On construit alors des horloges qui permettent de mesurer le temps.

2. Systèmes de coordonnées planes

Pour résoudre tout problème de physique et notamment de mécanique, il est nécessaire de repérer un point M dans l'espace. Il faut connaître les composantes du vecteur-position $\vec{OM}$ et des vecteurs qui pourront être définis à partir de celui-là.

¹ On rappelle qu'un élément chimique est défini par son nombre de charges c'est-à-dire son nombre d'électrons ou de protons : il y a autant de protons que d'électrons pour assurer la neutralité électrique de l'atome. Les isotopes d'un élément chimique sont les différents nucléides ou espèces de noyaux ; ils ont le même nombre de charges mais diffèrent par leur nombre de masse. Ce dernier correspondant au nombre de protons et de neutrons, les isotopes d'un élément chimique diffèrent par leur nombre de neutrons.

On choisit donc une base de projection dans laquelle on cherche à déterminer les composantes des vecteurs. Ce choix sera orienté par la géométrie du problème. Ce paragraphe présente les différentes solutions qui peuvent être envisagées pour repérer un point dans un plan. On se limitera en effet aux mouvements plans durant la première période et on verra en seconde période comment repérer un point dans l'espace.

2.1 Une base fixe : les coordonnées cartésiennes

a) Définition

La base de projection est définie par deux vecteurs unitaires formant une base directe. Le premier est choisi *a priori* et l'autre s'en déduit par la recherche du caractère direct de la base. Quelle que soit la position du point M considéré, la base est fixe et ne change pas.

En notant $\vec{u}_x$ et $\vec{u}_y$ les vecteurs unitaires sur chacun des axes de la base, on obtient :

$$\vec{OM} = x\vec{u}_x + y\vec{u}_y$$

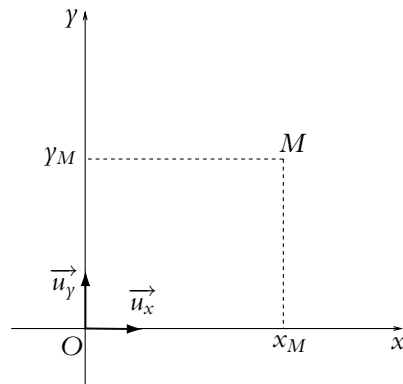


Figure 4.1 Coordonnées cartésiennes.

Ce système de coordonnées est celui auquel on pense le plus souvent mais ce n'est pas forcément le plus adapté à la symétrie du problème, notamment lorsque les points matériels se déplacent sur des cercles.

Pour éviter des confusions, les coordonnées cartésiennes de M seront notées x_M et y_M sur les figures, mais x et y dans le texte pour ne pas l'alourdir.

b) Déplacement élémentaire

Le déplacement élémentaire s'obtient en faisant varier de manière élémentaire chacune des coordonnées du point M .

Dans le cas des coordonnées cartésiennes, le déplacement élémentaire d'un point M de coordonnées (x, y) correspond à son déplacement jusqu'au point M' de coordonnées $(x+dx, y+dy)$. On a donc :

$$\begin{aligned} d\vec{OM} &= \vec{OM'} - \vec{OM} = \vec{MM'} \\ &= dx\vec{u}_x + dy\vec{u}_y \end{aligned}$$

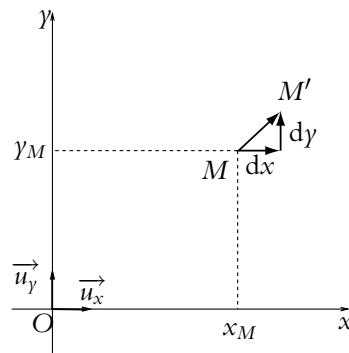


Figure 4.2 Déplacement élémentaire en coordonnées cartésiennes planes.

Ce résultat peut s'obtenir géométriquement en sommant les déplacements élémentaires liés à la variation d'une seule variable, les autres restant constantes.

Ceci est possible car les deux variables sont indépendantes et la base fixe.

Ainsi lorsque y (respectivement x) reste fixe, le déplacement élémentaire se limite à $d\vec{OM} = dx\vec{u}_x$ (respectivement à $d\vec{OM} = dy\vec{u}_y$).

2.2 Une base mobile : les coordonnées polaires

a) Définition

La symétrie polaire consiste à privilégier un point O fixe autour duquel tourne le point M .

La position de M est alors définie par la distance r du point M au point O . On utilise en outre un angle orienté de rotation appelé *angle polaire* et noté habituellement θ . Cet angle de rotation est défini par rapport à un axe passant par O choisi arbitrairement.

Les coordonnées polaires de M sont donc (r, θ) .

Par définition, r est une distance et est donc positif. On peut noter alors que pour décrire la totalité des points du plan, l'angle θ doit décrire un segment d'amplitude 2π . Habituellement on prend θ dans l'intervalle $[0, 2\pi]$.

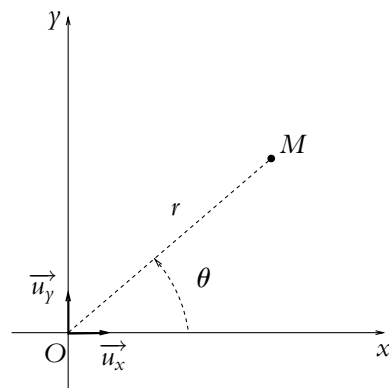


Figure 4.3 Coordonnées polaires.

b) Lien avec les coordonnées cartésiennes

On peut relier les coordonnées polaires d'un point à ses coordonnées cartésiennes. Il suffit pour cela de choisir l'axe (Ox) par exemple comme axe de référence pour l'angle de rotation θ et de projeter sur les axes (Ox) et (Oy) :

$$\begin{cases} x = r \cos \theta \\ y = r \sin \theta \end{cases}$$

On peut inverser ces relations pour obtenir l'expression des coordonnées cylindriques en fonction des coordonnées cartésiennes :

$$\begin{cases} r = \sqrt{x^2 + y^2} \\ \theta = \text{Arctan} \frac{y}{x} [\pi] \end{cases}$$

La valeur de θ sera précisée dans l'intervalle $[0, 2\pi]$ en fonction des signes respectifs de x et y qui sont ceux de $\cos \theta$ et $\sin \theta$.

Ceci n'est possible qu'en dehors de l'origine O . En ce point, $x = y = 0$, donc la tangente de l'angle θ en O est une forme indéterminée. On notera que l'impossibilité à exprimer la tangente de θ à l'origine ne porte pas à conséquence : en ce point, l'angle polaire n'est pas défini.

c) Base locale polaire

Les bases locales correspondent à des bases définies en chaque point M de l'espace et adaptées au type de coordonnées choisies.

On va définir ici la base locale des coordonnées polaires. Il s'agit de remplacer les vecteurs de base $(\vec{u}_x, \vec{u}_y)$ adaptés aux coordonnées cartésiennes par des vecteurs de base adaptés aux coordonnées polaires.

On construit ainsi la base locale définie par les deux vecteurs suivants :

- le premier vecteur noté $\vec{u}_r$ est un vecteur unitaire dont la direction et le sens sont ceux du vecteur $\vec{OM}$,
- le second noté $\vec{u}_\theta$ est un vecteur unitaire obtenu par rotation d'un angle $+\frac{\pi}{2}$ dans le plan à partir de $\vec{u}_r$.

Les vecteurs $\vec{u}_r$ et $\vec{u}_\theta$ dépendent du point M : la base est donc mobile. Les vecteurs de base se déplacent en même temps que le point M , d'où la dénomination « locale » donnée à cette base.

Le vecteur-position $\vec{OM}$ s'écrit dans cette base :

$$\vec{OM} = r \vec{u}_r.$$

On peut exprimer les vecteurs de cette base locale dans la base des coordonnées cartésiennes : il suffit de projeter les vecteurs $\vec{u}_r$ et $\vec{u}_\theta$ sur les axes (Ox) et (Oy) .

On obtient :

$$\begin{cases} \vec{u}_r = \cos \theta \vec{u}_x + \sin \theta \vec{u}_y \\ \vec{u}_\theta = -\sin \theta \vec{u}_x + \cos \theta \vec{u}_y \end{cases}$$

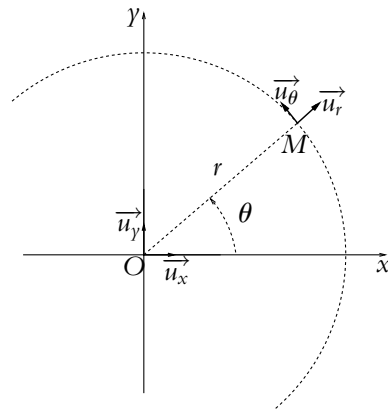


Figure 4.4 Base locale des coordonnées polaires.

Tout vecteur $\vec{w}$ peut se décomposer sur cette base selon :

$$\vec{w} = w_r \vec{u}_r + w_\theta \vec{u}_\theta$$

w_r est appelé *composante radiale* et w_θ *composante orthoradiale*.

d) Déplacement élémentaire

Le déplacement élémentaire d'un point M de coordonnées (r, θ) correspond à son déplacement jusqu'au point M' de coordonnées $(r + dr, \theta + d\theta)$. On a donc :

$$d\vec{OM} = \vec{OM'} - \vec{OM} = \vec{MM'}$$

Les variables r, θ sont indépendantes : l'expression du déplacement élémentaire est donc la somme des déplacements élémentaires correspondant à la variation d'une seule variable.

- une variation dr de r correspond à un mouvement le long de la droite (OM) soit $dr \vec{u}_r$,
- une variation $d\theta$ de θ correspond à un déplacement circulaire de centre O , de rayon r et d'angle $d\theta$ donc de longueur $rd\theta$ suivant $\vec{u}_\theta$, soit $rd\theta \vec{u}_\theta$,

Au final, le déplacement élémentaire s'écrit :

$$d\vec{OM} = dr \vec{u}_r + rd\theta \vec{u}_\theta$$

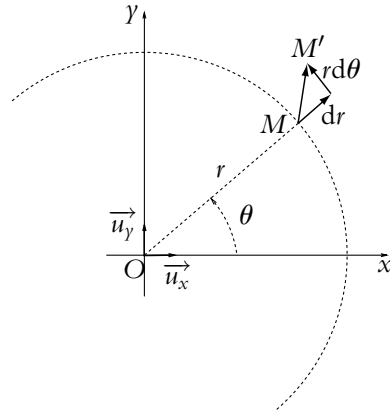


Figure 4.5 Déplacements élémentaires en coordonnées polaires.

e) Dérivées des vecteurs de base par rapport à θ

Quand $M(r, \theta)$ devient $M'(r + dr, \theta + d\theta)$, les vecteurs $\vec{u}_r$ et $\vec{u}_\theta$ deviennent respectivement $\vec{u}_r + d\vec{u}_r$ et $\vec{u}_\theta + d\vec{u}_\theta$.

Or les expressions de $\vec{u}_r$ et $\vec{u}_\theta$ établies dans la base fixe $(\vec{u}_x, \vec{u}_y)$ ne dépendent que de θ :

$$\begin{cases} \vec{u}_r = \cos \theta \vec{u}_x + \sin \theta \vec{u}_y \\ \vec{u}_\theta = -\sin \theta \vec{u}_x + \cos \theta \vec{u}_y \end{cases} \quad (4.1)$$

Les variations induites par le déplacement de M ne sont donc fonction que de θ . En effet, quand r varie de dr à θ constant, la direction du vecteur $\vec{OM}$ ne change pas, la base locale non plus.

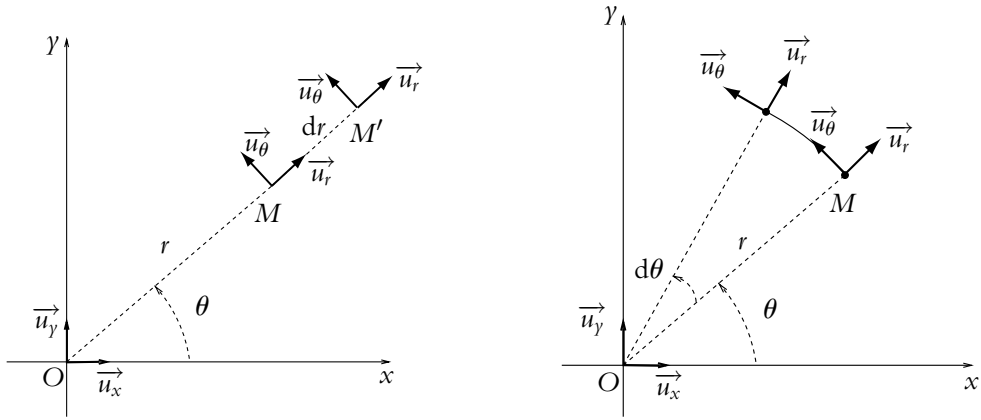


Figure 4.6 Déplacements élémentaires en coordonnées polaires.

Sur la figure de droite, l'angle $d\theta$ a été fortement exagéré pour la clarté du schéma, l'arc de cercle MM' est en réalité assimilable à un segment de droite.

En dérivant par rapport à θ les expressions précédentes, on déduit :

$$\begin{cases} \frac{d\vec{u}_r}{d\theta} = -\sin\theta\vec{u}_x + \cos\theta\vec{u}_y = \vec{u}_\theta \\ \frac{d\vec{u}_\theta}{d\theta} = -\cos\theta\vec{u}_x - \sin\theta\vec{u}_y = -\vec{u}_r \end{cases}$$

qu'on peut écrire :

$$\begin{cases} d\vec{u}_r = d\theta\vec{u}_\theta \\ d\vec{u}_\theta = -d\theta\vec{u}_r \end{cases}$$

L'expression du déplacement élémentaire s'obtient alors en différenciant le vecteur $\vec{OM} = r\vec{u}_r$:

$$d\vec{OM} = d(r\vec{u}_r) = dr\vec{u}_r + r d\vec{u}_r = dr\vec{u}_r + r d\theta\vec{u}_\theta$$

On obtient bien sûr le même résultat qu'en sommant les déplacements élémentaires.

3. Description cinématique du mouvement d'un point matériel

3.1 Référentiel et temps absolu

On appelle *référentiel* un système d'axes lié à un observateur, ce dernier étant muni d'une horloge pour mesurer le temps.

On se place dans le cadre de la mécanique classique, par opposition à la mécanique relativiste : le temps est absolu, c'est-à-dire que la mesure du temps est la même dans tout référentiel donc pour tout observateur.

3.2 Position, équation horaire et trajectoire

On définit la position d'un point matériel M dans un référentiel à l'aide du vecteur-position $\overrightarrow{OM}$ où O est un point fixe dans le référentiel. O peut être l'observateur mais ce n'est pas obligatoire ; la seule condition est que O soit fixe par rapport à ce dernier. On doit donc donner les composantes de $\overrightarrow{OM}$ dans la base de projection choisie, à savoir dans le système de coordonnées choisi.

En première période, on se limite à l'étude des mouvements plans : les deux composantes de $\overrightarrow{OM}$ dans la base de projection plane choisie, cartésienne ou polaire, suffisent.

Le mouvement du point matériel M conduit à étudier l'évolution de ce dernier au cours du temps et à donner l'équation horaire du mouvement par $\overrightarrow{OM}(t)$ ou ses composantes dans la base de projection choisie.

La courbe décrite par M au cours du mouvement est appelée *trajectoire*.

Il est important de noter que toutes ces définitions sont relatives au référentiel choisi. On verra en deuxième période comment on peut passer d'un référentiel à un autre.

3.3 Vitesse et hodographe

Soit M la position d'un point matériel à l'instant t et M' celle du même point matériel à l'instant $t + \delta t$.

On définit la *vitesse* de ce point à l'instant t par :

$$\vec{v}(M) = \lim_{\delta t \rightarrow 0} \frac{\overrightarrow{MM'}}{\delta t}$$

En introduisant un point fixe O par rapport au référentiel considéré, on obtient :

$$\overrightarrow{MM'} = \overrightarrow{MO} + \overrightarrow{OM'} = \overrightarrow{OM'} - \overrightarrow{OM} = d\overrightarrow{OM}$$

Alors la définition de la dérivée d'une fonction vectorielle permet d'écrire :

$$\vec{v}(M) = \frac{d\overrightarrow{OM}}{dt}$$

► **Remarque :** Une conséquence immédiate de cette définition est que le vecteur vitesse est tangent à la trajectoire au point considéré.

La vitesse dépend du référentiel dans lequel on la calcule. En effet, un point fixe par rapport à un référentiel $\mathcal{R}_1$ ne l'est plus dans un référentiel $\mathcal{R}_2$ en mouvement par rapport à $\mathcal{R}_1$.

En revanche, si on envisage un autre point fixe O' alors :

$$\frac{d\overrightarrow{O'M}}{dt} = \frac{d\overrightarrow{O'O}}{dt} + \frac{d\overrightarrow{OM}}{dt} = \frac{d\overrightarrow{OM}}{dt}$$

La vitesse d'un point M dans un référentiel $\mathcal{R}$ est indépendante du point fixe choisi dans ce référentiel.

Comme la position, la vitesse évolue au cours du temps. On obtient l'analogue de l'équation horaire (qui ne porte pas de nom) en considérant l'évolution du vecteur $\overrightarrow{v}(t)$ au cours du temps.

On peut également s'intéresser à l'allure de la courbe décrite par le vecteur vitesse en définissant l'*hodographe*. Il s'agit du lieu des points P tels qu'à chaque instant :

$$\overrightarrow{OP}(t) = \overrightarrow{v}(t)$$

3.4 Accélération

L'*accélération* est, par définition, la dérivée du vecteur-vitesse ou encore la dérivée seconde du vecteur-position par rapport au temps :

$$\overrightarrow{a}(M) = \frac{d\overrightarrow{v}(M)}{dt} = \frac{d^2\overrightarrow{OM}}{dt^2}$$

4. Expressions de la vitesse et de l'accélération

4.1 Coordonnées cartésiennes

On part du vecteur-position :

$$\overrightarrow{OM} = x(t)\overrightarrow{u_x} + y(t)\overrightarrow{u_y}$$

et on dérive par rapport au temps sachant que les vecteurs $\overrightarrow{u_x}$ et $\overrightarrow{u_y}$ sont constants au cours du temps.

Donc la vitesse s'écrit :

$$\overrightarrow{v}(M) = \frac{d\overrightarrow{OM}}{dt} = \dot{x}\overrightarrow{u_x} + \dot{y}\overrightarrow{u_y}$$

en notant $\dot{x} = \frac{dx}{dt}$ et $\dot{y} = \frac{dy}{dt}$.

De même pour l'accélération :

$$\overrightarrow{a}(M) = \frac{d\overrightarrow{v}(M)}{dt} = \ddot{x}\overrightarrow{u_x} + \ddot{y}\overrightarrow{u_y}$$

en notant $\ddot{x} = \frac{d^2x}{dt^2}$ et $\ddot{y} = \frac{d^2y}{dt^2}$.

4.2 Coordonnées polaires

Le vecteur-position s'écrit :

$$\overrightarrow{OM} = r(t) \vec{u}_r(t)$$



Le vecteur $\vec{u}_r$ évolue au cours du mouvement du point M .

On en déduit l'expression de la vitesse :

$$\vec{v}(M) = \frac{d\overrightarrow{OM}}{dt} = \frac{dr}{dt} \vec{u}_r + r \frac{d\vec{u}_r}{dt} = \frac{dr}{dt} \vec{u}_r + r \frac{d\vec{u}_r}{d\theta} \frac{d\theta}{dt}$$

soit en utilisant les résultats vus au paragraphe 2.2 et notant $\dot{r} = \frac{dr}{dt}$ et $\dot{\theta} = \frac{d\theta}{dt}$:

$$\vec{v}(M) = \dot{r} \vec{u}_r + r \dot{\theta} \vec{u}_\theta$$

On peut retrouver ces résultats géométriquement : il suffit de sommer les déplacements élémentaires correspondant aux variations de chacune des variables pendant un intervalle de temps dt . On se reportera aux résultats obtenus précédemment pour les expressions des déplacements élémentaires en coordonnées polaires.

Pour l'accélération, on dérive l'expression obtenue pour la vitesse :

$$\vec{a}(M) = \frac{d\vec{v}(M)}{dt} = \ddot{r} \vec{u}_r + \dot{r} \frac{d\vec{u}_r}{dt} + \dot{\theta} \vec{u}_\theta + r \frac{d\dot{\theta}}{dt} \vec{u}_\theta + r \dot{\theta} \frac{d\vec{u}_\theta}{dt}$$

soit en introduisant les dérivées des vecteurs unitaires par rapport à l'angle de rotation :

$$\begin{aligned} \vec{a}(M) &= \ddot{r} \vec{u}_r + \dot{r} \frac{d\vec{u}_r}{d\theta} \frac{d\theta}{dt} + \dot{\theta} \vec{u}_\theta + r \ddot{\theta} \vec{u}_\theta + r \dot{\theta} \frac{d\vec{u}_\theta}{d\theta} \frac{d\theta}{dt} \\ &= \ddot{r} \vec{u}_r + \dot{r} \vec{u}_\theta \dot{\theta} + \dot{\theta} \vec{u}_\theta + r \ddot{\theta} \vec{u}_\theta + r \dot{\theta} (-\vec{u}_r) \dot{\theta} \end{aligned}$$

et finalement :

$$\vec{a}(M) = (\ddot{r} - r\dot{\theta}^2) \vec{u}_r + (2\dot{r}\dot{\theta} + r\ddot{\theta}) \vec{u}_\theta$$

en notant $\ddot{r} = \frac{d^2r}{dt^2}$ et $\ddot{\theta} = \frac{d^2\theta}{dt^2}$.

4.3 Remarque : différence entre base de projection et référentiel

Les résultats qui viennent d'être obtenus permettent de souligner la différence entre base de projection (ou repère) et référentiel.

Le vecteur-vitesse en coordonnées polaires s'écrit $\vec{v}(M) = \dot{r}\vec{u}_r + r\dot{\theta}\vec{u}_\theta$. On obtient les composantes de la vitesse dans cette base de projection qui suit la particule dans son mouvement.

En revanche, il ne s'agit pas d'un référentiel mais d'une base de projection qui se déplace comme le point considéré.

La vitesse doit donc être définie relativement à un référentiel et non par rapport à une base de projection. Un référentiel est lié à un observateur muni d'un système d'axes. Une base de projection n'est qu'un outil permettant d'exprimer des équations vectorielles comme un ensemble d'équations scalaires plus faciles à manipuler. Elle n'est pas attachée à un observateur : ce n'est qu'un système d'axes dont les directions peuvent varier au cours du temps (Cf. base polaire). On pourra donc pour un même référentiel choisir une base de projection ou une autre suivant la commodité qu'elle apporte pour l'étude du mouvement. Il ne s'agit plus alors que d'aspects géométriques.

Par exemple, si on considère la rotation d'un point matériel sur un cercle, on choisira naturellement comme référentiel celui dans lequel le cercle est fixe. La base de projection pourra être soit la base fixe des coordonnées cartésiennes, soit la base mobile des coordonnées polaires.

Si on choisit comme référentiel d'étude celui qui suit le point matériel dans son mouvement, la vitesse du point matériel y est nulle. Ce dernier n'a donc pas grand intérêt pour l'étude de la trajectoire du point matériel.

5. Exemples de mouvements

5.1 Mouvements rectilignes

a) Définition

Un mouvement est dit *rectiligne* s'il correspond à un déplacement le long d'une droite fixe dans le référentiel d'étude. Dans ce cas, les vecteurs vitesse et accélération sont portés par la droite sur laquelle s'effectue le mouvement.

Le système de coordonnées le plus adapté est celui qui privilégie un axe qui sera la droite. Il est donc naturel de choisir les coordonnées cartésiennes.

b) Cas du mouvement rectiligne sinusoïdal

Si on choisit une base de coordonnées cartésiennes dont l'axe Ox est l'axe selon lequel s'effectue le mouvement, un mouvement rectiligne sera également sinusoïdal si l'équation horaire du mouvement est :

$$x(t) = X_0 \cos(\omega t + \varphi)$$

On notera que, dans ce cas, la vitesse et l'accélération sont également sinusoïdales :

$$\vec{v}(t) = v_x(t) \vec{u}_x = -X_0 \omega \sin(\omega t + \varphi) \vec{u}_x = X_0 \omega \cos\left(\omega t + \varphi + \frac{\pi}{2}\right) \vec{u}_x$$

$$\vec{a}(t) = a_x(t) \vec{u}_x = -X_0 \omega^2 \cos(\omega t + \varphi) \vec{u}_x = X_0 \omega^2 \cos(\omega t + \varphi + \pi) \vec{u}_x$$

Position et vitesse ainsi que vitesse et accélération sont en quadrature de phase c'est-à-dire que $x(t)$ et $v_x(t)$ ainsi que $v_x(t)$ et $a_x(t)$ sont déphasés de $\frac{\pi}{2}$ tandis que position et accélération sont en opposition de phase à savoir que $x(t)$ et $a_x(t)$ sont déphasés de π . Ce sera par exemple, le cas d'un point matériel attaché à un ressort et contraint à se déplacer sur une droite.

5.2 Mouvements à accélération constante

Un mouvement est à accélération constante si $\vec{a} = \vec{A}$ où $\vec{A}$ est un vecteur constant indépendant du temps.

L'expression de la vitesse est alors par intégration : $\vec{v}(t) = \vec{A}t + \vec{v}_0$ où $\vec{v}_0$ est la vitesse à l'instant $t = 0$.

On en déduit que :

- le mouvement est rectiligne si $\vec{A}$ et $\vec{v}_0$ sont colinéaires,
- sinon le mouvement s'effectue dans le plan défini par les vecteurs $\vec{A}$ et $\vec{v}_0$ passant par la position occupée par le point matériel à l'instant initial.

Ce sera par exemple, le cas d'un point matériel en chute libre en l'absence de tout frottement.

5.3 Mouvements uniforme, accéléré et décéléré

On distingue trois types de mouvements suivant l'évolution de la norme du vecteur vitesse :

- les mouvements uniformes :
un mouvement sera dit *uniforme* si la norme du vecteur vitesse est constante,
- les mouvements accélérés :
un mouvement sera dit *accéléré* si la norme de la vitesse augmente,
- les mouvements décélérés :
un mouvement sera dit *décéléré* si la norme de la vitesse diminue.

On peut interpréter ces définitions à partir de l'orientation relative des vecteurs vitesse et accélération. En effet, on a :

$$\frac{dv^2}{dt} = \frac{d\|\vec{v}(M)\|^2}{dt} = 2\vec{v}(M) \cdot \frac{d\vec{v}(M)}{dt} = 2\vec{v}(M) \cdot \vec{a}(M) = 2v(M)a(M)\cos\alpha$$

en notant α l'angle entre les vecteurs vitesse $\vec{v}(M)$ et accélération $\vec{a}(M)$.

On suppose qu'il y a mouvement donc que $v(M) \neq 0$. On en déduit que :

- pour les mouvements uniformes $a(M) = 0$ ou $\alpha = \pm \frac{\pi}{2}$ (c'est-à-dire que les vecteurs vitesse et accélération sont perpendiculaires),
- pour les mouvements accélérés $\cos \alpha > 0$ soit $\alpha \in [0, \frac{\pi}{2}]$,
- pour les mouvements décélérés $\cos \alpha < 0$ soit $\alpha \in [\frac{\pi}{2}, \pi]$.

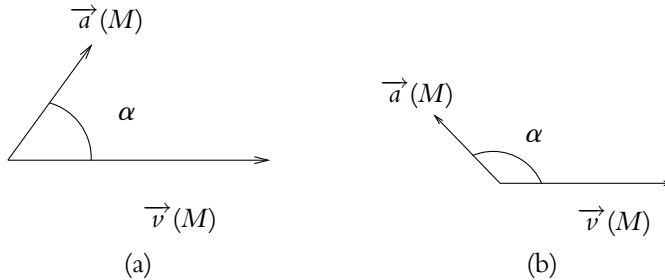


Figure 4.7 Positions relatives des vecteurs vitesse et accélération pour les mouvements accélérés (a) et décélérés (b).

Dans le cas d'un mouvement rectiligne où les vecteurs vitesse et accélération sont colinéaires ($\alpha = 0$), on a :

- un mouvement uniforme si $a = 0$,
- un mouvement accéléré si $\vec{v}(M)$ et $\vec{a}(M)$ sont de même sens,
- un mouvement décéléré si $\vec{v}(M)$ et $\vec{a}(M)$ sont de sens opposé.

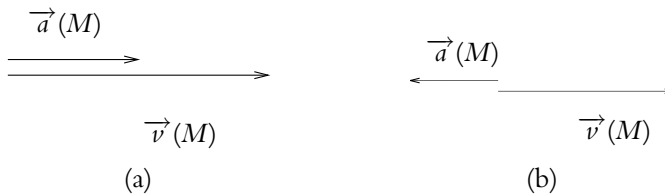


Figure 4.8 Mouvements rectilignes accéléré (a) et décéléré (b).

Dans le cas où $\frac{dv^2}{dt}$ est constant et non nul (ce qui correspond également à la constante du produit scalaire des vecteurs vitesse et accélération), on parle de mouvement uniformément accéléré ou uniformément décéléré.

On notera que dans le cas d'un mouvement rectiligne, il est uniforme si le vecteur accélération est nul puisque ce vecteur doit être à la fois colinéaire au vecteur vitesse du fait du caractère rectiligne et perpendiculaire à ce même vecteur vitesse du fait du caractère uniforme.

5.4 Mouvement circulaire

a) Définition

Un mouvement est dit *circulaire* si le point M se déplace sur un cercle fixe de centre O et de rayon R . La trajectoire est incluse dans le cercle.

Les coordonnées les plus adaptées à la description du mouvement sont naturellement les coordonnées polaires. Dans cette base de projection, les équations horaires du mouvement sont :

$$\begin{cases} r = R = \text{constante} \\ \theta(t) = (Ox, \overrightarrow{OM}) \end{cases}$$

b) Expression de la vitesse

Dans ce cas, on a :

$$\overrightarrow{OM} = R\overrightarrow{u_r}(t)$$

En dérivant, on obtient le vecteur vitesse :

$$\overrightarrow{v}(M) = R\dot{\theta}\overrightarrow{u_\theta}$$

On notera que le vecteur vitesse est orthoradial.

On appelle *vitesse angulaire* la quantité : $\omega = \dot{\theta}$. Le module de la vitesse s'écrit alors : $v = R\omega$.

Dans le cas d'un mouvement uniforme, la vitesse angulaire ω est constante puisque v et R sont constants.

c) Expression de l'accélération

Le vecteur vitesse s'écrit :

$$\overrightarrow{v}(M) = R\dot{\theta}(t)\overrightarrow{u_\theta}(t)$$

En suivant la même démarche que pour obtenir l'expression de l'accélération en coordonnées polaires, on en déduit, par dérivation par rapport au temps, l'expression de l'accélération (on rappelle que R est constant) :

$$\overrightarrow{a}(M) = \frac{d\overrightarrow{v}(M)}{dt} = R\ddot{\theta}\overrightarrow{u_\theta} + R\dot{\theta}\frac{d\overrightarrow{u_\theta}}{dt} = R\ddot{\theta}\overrightarrow{u_\theta} + R\dot{\theta}\frac{d\overrightarrow{u_\theta}}{d\theta}\frac{d\theta}{dt} = R\ddot{\theta}\overrightarrow{u_\theta} - R\dot{\theta}^2\overrightarrow{u_r}$$

On peut noter que dans le cas d'un mouvement circulaire uniforme les vecteurs $\overrightarrow{a}(M)$ et $\overrightarrow{v}(M)$ sont orthogonaux, donc :

$$\frac{d\omega}{dt} = 0$$

La vitesse angulaire $\omega = \dot{\theta}$ est constante.

► **Remarque :** Par conséquent, l'accélération est uniquement suivant $\overrightarrow{u_r}$. On retiendra qu'elle n'est pas nulle car le vecteur vitesse change de direction et n'est donc pas constant.

A. Applications directes du cours

1. Test d'accélération d'une voiture

Une voiture est chronométrée pour un test d'accélération avec départ arrêté (vitesse initiale nulle).

1. Elle est chronométrée à 26,6 s au bout d'une distance $D = 180$ m. Déterminer l'accélération (supposée constante) et la vitesse atteinte à la distance D .
2. Quelle est alors la distance d'arrêt pour une décélération de 7 m.s^{-2} ?

2. Interpellation pour vitesse excessive

Un conducteur roule à vitesse constante v_0 sur une route rectiligne. Comme il est en excès de vitesse à 110 km.h^{-1} , un gendarme à moto démarre à l'instant où la voiture passe à sa hauteur et accélère uniformément. Le gendarme atteint la vitesse de 90 km.h^{-1} au bout de 10 s.

1. Quel sera le temps nécessaire au motard pour rattraper la voiture ?
2. Quelle distance aura-t-il parcourue ?
3. Quelle vitesse aura-t-il alors atteinte ?

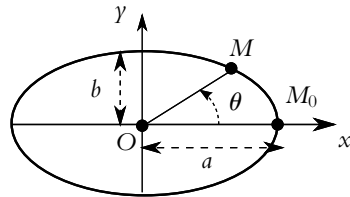
3. Satellite géostationnaire

Un satellite géostationnaire est en mouvement circulaire uniforme autour de la Terre. Il ressent une accélération $a = g_0 \left(\frac{R}{r} \right)^2$, où $R = 6400 \text{ km}$ est le rayon de la terre, $g_0 = 9,8 \text{ m.s}^{-2}$ et r le rayon de l'orbite. La période de révolution du satellite est égale à la période de rotation de la terre sur elle-même.

1. Calculer la période T de rotation de la Terre en secondes, puis la vitesse angulaire Ω .
2. Déterminer l'altitude de l'orbite géostationnaire.

4. Mouvement sur une ellipse

Un point se déplace sur une ellipse d'équation cartésienne $\left(\frac{x}{a} \right)^2 + \left(\frac{y}{b} \right)^2 = 1$. On note θ l'angle que fait $\overrightarrow{OM}$ avec l'axe Ox .



1. Les coordonnées d'un point M sur l'ellipse peuvent s'écrire $x(t) = \alpha \cos(\omega t + \phi)$ et $y(t) = \beta \sin(\omega t + \psi)$ où l'on suppose que ω est une constante.

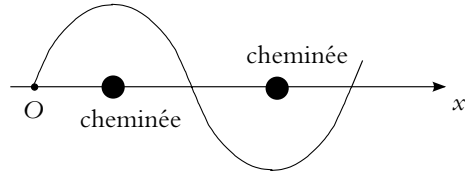
- a) À $t = 0$, le mobile est en M_0 . Déterminer α , ϕ et ψ .
- b) Des autres données, déduire β .

2. a) Déterminer les composantes de la vitesse $(\dot{x}, \dot{y})$ et de l'accélération $(\ddot{x}, \ddot{y})$ dans la base cartésienne.

b) Montrer que l'accélération est de la forme $\vec{a} = -k\overrightarrow{OM}$ où k est à déterminer.

5. Slalom entre des cheminées ; Star Wars, épisode 2

Dans cet épisode de la Guerre des Étoiles, on peut assister à une course poursuite de « speeder » entre des cheminées d'usine. On suppose que le véhicule suit une trajectoire sinusoïdale de slalom entre les cheminées alignées selon l'axe Ox . Elles sont espacées d'une distance $L = 200$ m.



1. Le véhicule conserve une vitesse v_0 constante selon Ox . Il met un temps $t_t = 12$ s pour revenir sur l'axe après la sixième cheminée. En déduire la vitesse v_0 . Effectuer l'application numérique.
2. Déterminer l'amplitude de la sinusoïde pour que l'accélération reste inférieure à $10g$ en valeur absolue, avec $g = 9,8 \text{ m.s}^{-2}$. Que penser des valeurs obtenues ?

B. Exercices et problèmes

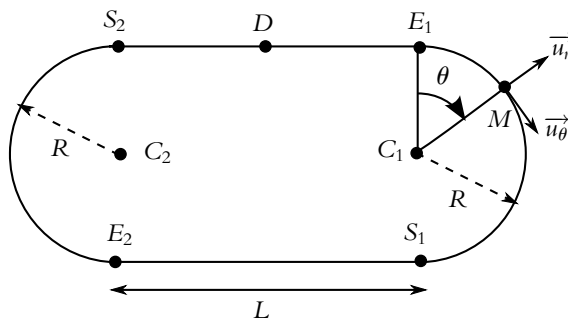
1. Courses entre véhicules radio-commandés

Deux modèles réduits de voitures radio-commandées ont des performances différentes : le premier a une accélération de $4,0 \text{ m.s}^{-1}$, le second de $5,0 \text{ m.s}^{-1}$. Cependant l'utilisateur de la première voiture a plus de réflexes que celui de la seconde, ce qui lui permet de la faire démarrer $1,0$ s avant le second.

1. Déterminer le temps nécessaire au véhicule à plus grande accélération pour rattraper l'autre.
2. Les deux modèles réduits participent à des courses de 100 m et 200 m. Est-il possible que le perdant du 100 m prenne sa revanche au 200 m ?
3. Calculer pour les deux courses la vitesse finale de chacun des véhicules.

2. Parcours d'un cycliste sur un vélodrome

On s'intéresse à un cycliste, considéré comme un point matériel M , qui s'entraîne sur un vélodrome constitué de deux demi-cercles reliés par deux lignes droites (figure ci-dessous). Données : $L = 62$ m et $R = 20$ m. Le cycliste part de D avec une vitesse nulle.



1. Il exerce un effort constant ce qui se traduit par une accélération constante a_1 jusqu'à l'entrée E_1 du premier virage. Calculer le temps t_{E_1} de passage en E_1 ainsi que la vitesse V_{E_1} en fonction de a_1 et L .
2. Dans le premier virage, le cycliste a une accélération tangentielle (suivant $\vec{u}_\theta$) constante égale a_1 . Déterminer le temps t_{S_1} de passage en S_1 ainsi que la vitesse v_{S_1} en fonction de a_1 , L et R .
3. a) De même, en considérant l'accélération tangentielle constante tout au long du premier tour et égale à a_1 , déterminer les temps t_{E_2} , t_{S_2} et t_D (après un tour), ainsi que les vitesses correspondantes.
b) La course s'effectue sur quatre tours (1 km) mais on ne s'intéresse donc qu'au premier effectué respectivement en $T_1 = 18,155$ s (Temps du britannique Chris Hoy aux Championnats du monde de 2007). Déterminer la valeur de l'accélération a_1 ainsi que la vitesse atteinte en D . La vitesse mesurée sur piste est d'environ 60 km.h^{-1} que doit-on modifier dans le modèle pour se rapprocher de la réalité?

3. La course poursuite

Quatre robots repérés par les points A , B , C et D sont initialement aux sommets d'un carré de côté a , de centre O . Le robot A est programmé pour toujours se diriger vers B avec une vitesse de norme constante v , de même B se dirige toujours vers C à vitesse v , C vers D et D vers A avec la même vitesse. On s'intéresse au mouvement de A (les trois autres s'en déduisent par rotation).

1. En utilisant des coordonnées polaires de centre O , déterminer les équations paramétriques $r(t)$ et $\theta(t)$ du mouvement de A en fonction de t , v .
2. Déterminer l'instant d'arrivée en O .

4. Mouvement de l'extrémité d'une barre, d'après ENAC 2003

Une barre rectiligne AB de longueur $2b$ se déplace dans le référentiel $\mathcal{R}$ repéré par $(O, \vec{u}_x, \vec{u}_y, \vec{u}_z)$ de telle sorte que (figures 4.9 et 4.10) :

- son extrémité A se trouve sur le demi-axe positif Oz ;
- son extrémité B décrit le demi-cercle du plan (xOy) de centre $I(0, b, 0)$ et de rayon b , à la vitesse angulaire ω constante et positive. À l'instant $t = 0$, B se trouve en O .

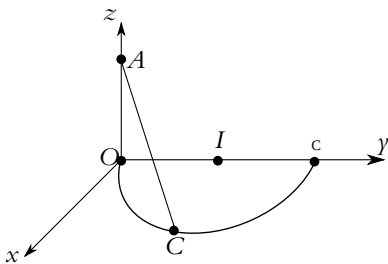


Figure 4.9

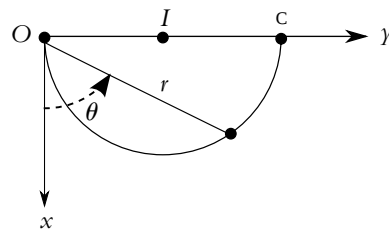


Figure 4.10

L'exercice ne nécessite aucune connaissance de mécanique du solide.

1. Déterminer la durée T du mouvement.
2. a) On note ϕ l'angle $(\overrightarrow{OI}, \overrightarrow{OB})$. Déterminer une relation simple entre ϕ et θ .
b) Établir les expressions en fonction du temps t des coordonnées polaires r et θ de B (figure 4.10).
3. a) Déterminer l'angle $\alpha = (\overrightarrow{AO}, \overrightarrow{AB})$ en fonction de ω et t .
b) Décrire le mouvement de la barre entre l'instant initial et l'instant final.
4. Calculer les coordonnées cartésiennes X , Y et Z du milieu J de la barre.
5. a) Déterminer la vitesse $\overrightarrow{v}$ et l'accélération $\overrightarrow{a}$ de J , ainsi que leurs normes.

5

Principes de la dynamique newtonienne

Dans le chapitre précédent, on ne s'est intéressé qu'aux aspects cinématiques à savoir l'analyse du mouvement et la description des liens entre les divers vecteurs décrivant ce dernier (position, vitesse, accélération). On ne s'est jamais posé la question de déterminer les causes et l'origine du mouvement. C'est ce point qui va maintenant être envisagé. Cela conduira à introduire la notion de force ainsi qu'à énoncer les trois lois de Newton. Avant cela, il est nécessaire de définir des éléments cinétiques du système dont on n'avait aucune utilité en cinématique.

1. Éléments cinétiques d'un point matériel

1.1 Masse

D'un point de vue cinématique, un point matériel est décrit à l'aide des vecteurs position, vitesse et accélération. Cependant quelques exemples de la vie courante mettent en évidence que le comportement d'un corps ne dépend pas uniquement de ces paramètres cinématiques. Ainsi un joueur de ping-pong sera dans l'impossibilité de renvoyer la balle si celle-ci est remplacée par une boule de billard. De même, si un enfant joue avec des boules de pétanque, il ne pourra pas lancer celles-ci très loin. Il est donc nécessaire d'introduire une grandeur physique mesurant la capacité du corps à résister au mouvement qu'on souhaite lui imposer. Cette propriété correspond à l'inertie du système.

La grandeur introduite est fondamentale au niveau dynamique et s'appelle *masse inerte* ou *masse inertielle* du corps. Il s'agit d'un scalaire positif qui est d'autant plus grand que le corps s'oppose au mouvement. Son unité légale est le kilogramme, de symbole kg. On constate expérimentalement que cette grandeur est proportionnelle à la quantité de matière composant le corps. C'est donc une grandeur additive, c'est-à-dire que la masse de l'ensemble formé par deux corps est égale à la somme des masses de chacun d'entre eux.

La mesure de cette grandeur s'obtient en pratique à l'aide d'une balance en utilisant la mesure du poids du corps : $\vec{p} = m\vec{g}$ où m est alors la *masse pesante*. Plusieurs problèmes apparaissent avec cette méthode. Tout d'abord, on verra que le champ de pesanteur $\vec{g}$ varie à la surface de la Terre. Il serait donc *a priori* nécessaire d'effectuer les réglages à chaque déplacement de la balance grâce à une référence. La référence universelle est un cylindre de platine iridié qui constitue un étalon conservé au Pavillon de Breteuil au sein du Bureau International des Poids et Mesure ; il s'agit de la seule référence qui est encore définie de nos jours par un étalon. Le deuxième point plus crucial encore est l'hypothèse selon laquelle masse inertielle et masse pesante (celle qui apparaît dans l'expression de l'interaction gravitationnelle et donc dans celle du poids) ne sont qu'une seule et même grandeur. On reviendra ultérieurement sur ce point et on admet pour l'instant l'égalité des masses inertielle et pesante.

Il s'agit d'une grandeur intrinsèque au point matériel et indépendante du référentiel dans lequel on se place.

1.2 Quantité de mouvement

La *quantité de mouvement* est une grandeur introduite par Isaac Newton pour formuler les lois de la mécanique portant son nom. Elle est définie par le vecteur :

$$\vec{p}_{/\mathcal{R}}(M) = m \vec{v}_{/\mathcal{R}}(M)$$

pour un point de masse m et de vecteur vitesse $\vec{v}_{/\mathcal{R}}(M)$ par rapport à un référentiel $\mathcal{R}$. Contrairement à la masse, cette quantité dépend du référentiel dans lequel on travaille puisqu'elle est fonction de la vitesse dans ce référentiel.

2. Notion de force

On s'intéresse dans ce paragraphe aux causes et à l'origine des mouvements et/ou de leur modification, c'est-à-dire aux interactions mécaniques entre le système et le milieu extérieur.

2.1 Point isolé

On dit qu'un point matériel est *isolé* lorsqu'il n'est soumis à aucune interaction mécanique avec l'extérieur.

Cette définition a-t-elle une réalité physique ? Il est en effet difficilement imaginable de soustraire un point matériel à l'influence de tout ce qui l'entoure. Le caractère isolé n'est donc qu'un modèle utilisé en mécanique. Dans le cadre de ce dernier, un point matériel sera considéré comme isolé lorsque toutes les interactions auxquelles il est soumis sont négligeables à l'échelle envisagée.

Dans la suite du paragraphe, on considère un système non isolé.

2.2 Définition de la notion de force

Un système peut être mis en mouvement ou, s'il est déjà animé d'un mouvement, ce dernier peut être modifié. Dans ce cas, il convient de rechercher les causes de la « modification » de l'état ou de la trajectoire d'un système. Cela conduit à définir la notion de forces :

On appelle *force* la grandeur vectorielle décrivant l'interaction capable de modifier et/ou de produire un mouvement ou une déformation du système.

La force est décrite par un vecteur ; le caractère vectoriel de la force apparaît dans des expériences simples. Par exemple, lorsqu'on tire sur un ressort, on constate que ce dernier se déplace le long d'une direction et dans un sens qui sont ceux de l'effort ou de la force qu'on exerce sur lui. Son allongement est d'autant plus grand que la force est importante. Trois paramètres : direction, sens et intensité interviennent pour déterminer l'action exercée sur le ressort. Un vecteur défini par ces trois mêmes paramètres décrit la force. Il faudra donc donner la direction, le sens et la norme d'une force pour la connaître parfaitement.

On distingue deux grandes catégories de forces :

- les forces à distance,
- les forces de contact.

On va donc aborder successivement ces deux types d'interactions.

2.3 Forces à distance

Actuellement on considère quatre types d'interaction à distance entre particules.

a) Interaction électromagnétique

Il s'agit de l'interaction entre particules chargées électriquement.

Cette interaction est décrite dans le cadre des équations de Maxwell qui sont à la base de tout l'électromagnétisme. On ne va pas s'intéresser ici à ces équations qui seront étudiées dans le cours de seconde année. La détermination des champs électriques et magnétiques dans le cas de distributions de charges et de courants permanents sera vue dans la partie consacrée à l'électromagnétisme statique du programme de première année.

Ce qui importe ici est l'influence mécanique d'un champ électrique et d'un champ magnétique sur une particule de charge q sans tenir compte de l'origine des champs. Dans un référentiel $\mathcal{R}$, la force subie par une particule de charge q et de vitesse $\vec{v}/\mathcal{R}$ soumise à un champ électrique $\vec{E}$ et à un champ magnétique $\vec{B}$ est décrite par la

force de Lorentz :

$$\vec{F} = q \left(\vec{E} + \vec{v}_{/R} \wedge \vec{B} \right)$$

Cette force traduit les interactions entre particules chargées dans la mesure où les champs électrique $\vec{E}$ et magnétique $\vec{B}$ sont créés par des particules chargées, immobiles ou en mouvement, ces champs dépendant *a priori* du point et du référentiel.

Dans le cas particulier de l'électrostatique où on ne s'intéresse qu'à des particules chargées immobiles, l'interaction prend la forme plus simple de la loi de Coulomb. La force subie par une charge q_2 placée en M_2 du fait de la présence en M_1 d'une charge q_1 s'écrit :

$$\overrightarrow{F_{1 \rightarrow 2}} = \frac{1}{4\pi\epsilon_0} \frac{q_1 q_2}{(M_1 M_2)^3} \overrightarrow{M_1 M_2}$$

en notant $\epsilon_0 = 8,854 \cdot 10^{-12} \text{ F.m}^{-1}$ la permittivité diélectrique du vide. On retrouve la formulation de la force de Lorentz en notant que :

$$\overrightarrow{F_{1 \rightarrow 2}} = q_2 \left(\frac{1}{4\pi\epsilon_0} \frac{q_1}{(M_1 M_2)^3} \overrightarrow{M_1 M_2} \right) = q_2 \overrightarrow{E_{1 \rightarrow 2}}$$

Dans cette expression, on ne fait que préciser l'expression du champ électrique créé par la charge q_1 placée en M_1 à l'endroit où se trouve la charge q_2 . C'est cette expression qui permet de définir la notion de champ électrique.

On constate dans le cas de la loi de Coulomb la portée infinie de l'interaction électrostatique. Il s'agit d'un résultat général pour l'ensemble des interactions électromagnétiques. Ce type d'interaction est particulièrement important car il est à l'origine d'un très grand nombre de phénomènes courants à notre échelle (à l'exclusion des phénomènes gravitationnels). En effet, outre les phénomènes purement électriques ou magnétiques (dont le chapitre sur le mouvement des particules chargées donnera quelques exemples), l'interaction électromagnétique est responsable des liaisons chimiques entre atomes d'une molécule ou d'un ion, des interactions entre constituants de la matière (atomes, molécules, ions) qui en assurent l'organisation et la cohésion ou encore des ondes électromagnétiques qui sont à l'origine des phénomènes optiques, des transmissions radios, etc.

b) Interaction gravitationnelle

Il s'agit d'une autre interaction à portée infinie qui intervient cette fois entre toutes particules possédant une masse dite gravitationnelle. Elle est décrite par la loi d'attraction universelle de Newton :

Tout point matériel de masse m_1 attire tout point matériel de masse m_2 avec une force dirigée le long de la droite qui les relie. Cette force varie comme l'inverse du

carré de la distance entre les particules et proportionnellement au produit de leurs masses.

La notion de masse gravitationnelle est définie à partir de cette loi d'attraction universelle : on constate expérimentalement que la force d'interaction est d'autant plus importante que la quantité de matière du système est grande. Il est donc nécessaire d'introduire une grandeur scalaire mesurant la quantité de matière du système, cette grandeur est appelée masse gravitationnelle.

L'interaction gravitationnelle se traduit donc par la force s'exerçant sur la masse m_2 située au point M_2 du fait de la présence de la masse m_1 placée en M_1 :

$$\overrightarrow{F_{1 \rightarrow 2}} = -G \frac{m_1 m_2}{(M_1 M_2)^3} \overrightarrow{M_1 M_2}$$

en notant $G = 6,67 \cdot 10^{-11} \text{ N.m}^2.\text{kg}^{-2}$ la constante de gravitation. On peut noter l'analogie entre les expressions de l'interaction gravitationnelle et de l'interaction électrostatique. La différence essentielle provient du fait que l'interaction électrostatique peut être attractive ou répulsive alors que l'interaction gravitationnelle n'est qu'attractive. En effet, les charges électriques peuvent être de même signe ou de signe contraire tandis que les masses sont toujours de même signe, à savoir positives.

c) Interaction nucléaire forte

Il s'agit de l'interaction qui assure la cohésion des noyaux par des forces entre neutrons et protons. Elle concerne donc les particules élémentaires. Elle a été imaginée pour expliquer l'existence de noyaux stables. Elle s'applique aux quarks et aux particules formées de quarks de la famille des mésons (méson K, méson π ...), des baryons (proton, neutron...) ou des hypérons (particule Λ , σ ...). Elle n'affecte pas la famille des leptons dont fait partie l'électron.

Sa portée est faible, de l'ordre du fermi ($1 \text{ fermi} = 10^{-15} \text{ m}$). Elle ne s'exerce qu'au sein du noyau ou au sein des étoiles à neutrons qui sont suffisamment compactes pour que les particules soient assez proches pour subir ce type d'interaction.

Il s'agit d'une interaction cent fois plus intense que l'interaction électromagnétique, ce qui lui a donné son nom.

d) Interaction nucléaire faible

Elle est d'une intensité de l'ordre de 10^{-5} fois plus faible que l'interaction forte, ce qui explique la dénomination employée.

Elle est également de très faible portée de l'ordre de 10^{-18} à 10^{-17} m et concerne les constituants des noyaux atomiques.

Elle est responsable notamment de la désintégration β d'un neutron en un proton, un électron et un neutrino.

e) Unification des interactions

Le souhait d'obtenir une théorie unique pour rendre compte de l'ensemble de ces quatre interactions est très ancien. En 1967, Samuel Weinberg et Abdus Salam (prix Nobel en 1979) ont proposé un modèle électrofaible unifiant les interactions faible et électromagnétique. Cette théorie a été vérifiée expérimentalement en 1983. À ce jour, aucune théorie unificatrice n'a été établie malgré de nombreuses recherches.

f) Quelles interactions considère-t-on ?

Les interactions forte et faible sont négligeables au niveau macroscopique. Dans la suite, on n'en tiendra pas compte puisqu'on se placera à cette échelle de taille. Il suffit ici de savoir que ces interactions existent et permettent d'expliquer des phénomènes incompréhensibles quand on ne considère que les interactions électromagnétique et gravitationnelle.

2.4 Forces de contact

Par opposition aux forces étudiées au paragraphe précédent qui avaient une action à distance, on s'intéresse maintenant aux actions qui n'existent que lors d'un contact entre systèmes.

a) Définition

Lorsqu'un point matériel n'est soumis qu'à des forces à distance, on dit qu'il est libre : sa trajectoire n'est pas astreinte à rester dans une zone plus ou moins confinée de l'espace. C'est le cas par exemple d'un corps qui tombe dans le champ de pesanteur en l'absence de tout frottement.

On s'intéresse maintenant aux forces qui n'interviennent qu'en présence d'un contact du point matériel avec un solide ou un fluide. On aura alors essentiellement des forces de liaison et des forces de frottement.

Il n'existe pas de théorie permettant de déterminer les forces de contact à l'aide de considérations microscopiques qu'on pourrait étendre à l'échelle macroscopique. Tout ce qui va suivre ne sera donc qu'une approche phénoménologique, c'est-à-dire obtenue expérimentalement, et valable uniquement dans le même contexte.

b) Tension d'un fil

Un fil de masse négligeable prend une forme rectiligne dès qu'il est tendu. On observe que la tension est la même en tout point du fil et que sa valeur dépend du mouvement du point. Dans le cas où le fil n'est plus tendu, la tension est nulle.

On modélise ces observations en considérant que la force exercée par un fil tendu sur un objet auquel il est attaché est équivalente à une tension $\vec{T}$ appliquée au point d'attache, dirigée le long du fil dans la direction de l'objet vers le fil et dont la norme dépend des autres forces appliquées.

c) Tension d'un ressort

Les ressorts envisagés dans la suite du cours et dans les exercices sont supposés parfaitement élastiques, c'est-à-dire qu'on peut négliger leur masse et qu'après une élongation ou une compression, ils reprennent leur forme et leur longueur initiales.

On constate alors que la tension qu'un ressort exerce sur un objet auquel il est accroché s'applique au point d'attache, le long du ressort, dans une direction opposée à son « allongement » (qui pourra être un réel allongement ou une compression) et que son intensité est proportionnelle à l'allongement.

On pourra donc adopter le modèle suivant :

$$\vec{F} = -k\Delta l \vec{u}_{ext}$$

où :

- k est une constante de proportionnalité caractéristique du ressort utilisé et appelée *raideur du ressort*, exprimée en N.m^{-1} dans les unités du système international,
- $\vec{u}_{ext}$ est un vecteur unitaire parallèle au ressort, toujours orienté vers l'extérieur du ressort,
- $\Delta l = l - l_0$ est l'allongement du ressort en notant respectivement l et l_0 la longueur du ressort et sa longueur à vide.

Si $\Delta l = l - l_0 > 0$ alors le ressort est étiré (ou tendu) ; si $\Delta l = l - l_0 < 0$ alors le ressort est comprimé.

La force qu'on vient de définir s'oppose à cette déformation du ressort par rapport à sa forme au repos (c'est-à-dire sans déformation) : c'est une force de rappel.

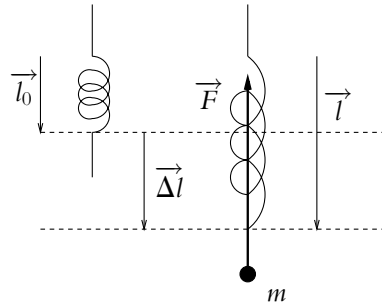


Figure 5.1 Définition de la tension d'un ressort.

d) Réaction d'un support solide

Un point matériel peut être astreint à se déplacer le long d'une surface ou d'une courbe. Son nombre de degrés de liberté ou nombre de paramètres nécessaires pour donner sa position est alors réduit et on parle de *point lié*. Le support exerce alors une force de contact.

Les actions du support peuvent être représentées par une force dite de réaction appliquée au point de contact entre le point matériel et le support solide. On peut la décomposer en la somme :

- d'une force $\vec{R}_N$ normale, au support.
- d'une force tangentielle dite force de frottement $\vec{R}_T$ appartenant au plan tangent au support.

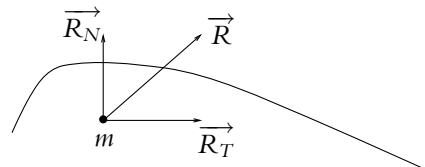


Figure 5.2 Définition de la réaction d'un support solide.

$$\vec{R} = \vec{R}_N + \vec{R}_T$$

Les propriétés des deux composantes $\vec{R}_N$ et $\vec{R}_T$ sont énoncées par les lois de Coulomb données ici à titre indicatif. On distingue deux cas :

- en l'absence de frottement, la réaction du support est normale à ce dernier :

$$\vec{R}_T = \vec{0}$$

et $\vec{R}_N$ est déterminée par les autres forces en présence,

- en présence d'un frottement, quand il y a mouvement, la force de frottement $\vec{R}_T$ a la même direction et est de sens opposé à la vitesse du point matériel. Sa norme est proportionnelle à celle de la réaction normale $\vec{R}_N$:

$$\|\vec{R}_T\| = f \|\vec{R}_N\|$$

où f est le coefficient de frottement dynamique.

Dans tout ce qui précède, il faut que le contact reste établi donc qu'il existe une réaction normale, ce qui se traduit par la condition suivante :

$$R_N > 0$$

Dès que la composante normale de la réaction s'annule ($R_N = 0$), le contact cesse.

L'origine physique de ces forces se situe au niveau des interactions électromagnétiques entre les molécules du support et celle du système.

e) Forces de frottement fluide

Lorsqu'un système se trouve dans un milieu fluide, il est souvent nécessaire de tenir compte de forces de frottement entre le système et le milieu dans lequel il évolue. On parle de frottement fluide.

La force traduisant cette interaction peut se décomposer en deux composantes :

- une force appelée *portance* qui est perpendiculaire à la vitesse du système par rapport au fluide ; c'est cette force qui assure aux avions la possibilité de se maintenir en altitude ;
- une force dite de *traînée*, $\vec{T}$, de sens opposé à la vitesse du système par rapport au fluide. On adopte généralement un modèle où la norme de cette force est proportionnelle à la vitesse du système par rapport au fluide lorsque celle-ci est faible et au carré de cette vitesse quand elle devient plus importante.

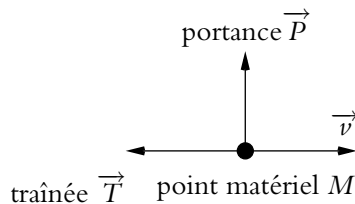


Figure 5.3 Définition de la portance et de la traînée.

Aux faibles vitesses, on écrit :

$$\vec{T} = -\lambda \vec{v}$$

où λ est une constante.

Aux plus fortes vitesses, l'expression devient :

$$\vec{T} = -\lambda' \nu \vec{v}$$

où λ' est une constante proportionnelle dans le cas d'un solide à un coefficient noté C_x , appelé *coefficient de traînée*, dépendant de la forme du solide, ainsi qu'à la masse volumique ρ du fluide dans lequel s'effectue le mouvement. À titre d'exemple, voici quelques valeurs classiques du coefficient de traînée, coefficient traduisant la capacité du système à pénétrer dans le fluide, ici l'air :

- pour une sphère lisse $C_x \simeq 0,4$,
- pour une voiture $C_x \simeq 0,3$,
- pour une balle de fusil $C_x \simeq 0,03$.

Aux vitesses intermédiaires ou aux vitesses très élevées, il n'y a plus de modèles simples.

Le critère définissant la limite entre vitesses faibles et grandes vitesses sera vu dans le cours de mécanique des fluides de seconde année (PC et PSI).

3. Les trois lois de Newton de la dynamique

Ces trois lois de Newton constituent la base de la dynamique newtonienne et de la mécanique classique (par opposition à la mécanique quantique ou relativiste). Il s'agit de principes au sens où ces lois sont postulées et non démontrées théoriquement. Elles sont donc considérées comme valables tant que l'expérience n'a pas contredit les conséquences qu'on peut déduire théoriquement de ces postulats. Ce sont justement les écarts entre certaines observations expérimentales et les prédictions théoriques issues des principes de Newton qui ont conduit à l'édification des théories de la mécanique quantique et de la mécanique relativiste en posant des principes différents de ceux de Newton. Pour les problèmes envisagés ici, elles sont bien évidemment valables.

3.1 Première loi de Newton : principe d'inertie

a) Énoncé du principe d'inertie

Il existe des référentiels privilégiés appelés *référentiels galiléens* dans lesquels un point matériel isolé est animé d'un mouvement rectiligne uniforme, c'est-à-dire que les vecteurs vitesse et quantité de mouvement sont constants au cours du temps.

b) Interprétation du principe d'inertie

Le principe d'inertie formule l'existence de référentiels particuliers : les référentiels galiléens dont il fournit une définition à partir du mouvement d'un point matériel isolé. On constate donc la différence essentielle apportée par la dynamique vis-à-vis de la cinématique : les référentiels ne jouent plus tous le même rôle.

Le principe d'inertie est un postulat initialement formulé par Galilée et repris par Newton. On reviendra en deuxième période sur la recherche d'un référentiel galiléen. Pour l'instant, on admet qu'il existe au moins un tel référentiel.

Ces référentiels sont tous en translation rectiligne uniforme les uns par rapport aux autres comme on le verra en seconde période.

3.2 Deuxième loi de Newton : principe fondamental de la dynamique

Ayant admis l'existence d'un référentiel galiléen au paragraphe précédent, on se place maintenant dans ce référentiel.

a) Énoncé du principe fondamental de la dynamique

Dans un référentiel galiléen, la dérivée par rapport au temps du vecteur quantité de mouvement est égale à la somme des forces s'exerçant sur le point matériel.

Mathématiquement, cela se traduit par la relation :

$$\frac{d\vec{p}_{/\mathcal{R}}(M)}{dt} = \sum_i \vec{f}_i$$

b) Interprétation

Ce principe relie donc un terme cinétique, la dérivée du vecteur quantité de mouvement, à un terme dynamique, la somme des forces traduisant les interactions subies par le point matériel.

Cette relation peut être utilisée dans les deux sens :

- obtenir la description cinématique du mouvement connaissant les forces subies par le point matériel,
- déterminer à partir de la connaissance du mouvement la somme des forces s'exerçant sur le point matériel.

Le principe fondamental de la dynamique est également appelé « relation fondamentale de la dynamique » ou « théorème de la quantité de mouvement ».

On préférera dans la suite l'expression de « principe fondamental de la dynamique » pour souligner qu'il s'agit d'une loi postulée, non démontrée mais vérifiée expérimentalement.

c) Première conséquence : cas des systèmes pseudo-isolés

Les systèmes *pseudo-isolés* sont définis comme les systèmes pour lesquels la somme des forces qui s'exercent sur eux est nulle :

$$\sum_i \vec{f}_i = \vec{0}$$

Pour de tels systèmes, le principe fondamental de la dynamique montre que la quantité de mouvement $\vec{p}_{/\mathcal{R}}(M)$ est constante au cours du temps. Ces systèmes sont donc animés d'un mouvement rectiligne uniforme dans un référentiel galiléen. Ceci explique leur dénomination : tout se passe comme s'ils étaient isolés.

d) Masse inertielle ou masse gravitationnelle

Deux notions de masse ont été introduites : la masse inertielle dans l'expression de la quantité de mouvement et du principe fondamental de la dynamique et la masse gravitationnelle pour l'interaction du même nom. Il n'y a *a priori* aucune raison que ces deux notions coïncident.

Les premières expériences pour établir l'équivalence entre masse inertielle m_i et masse gravitationnelle m_g sont dues à Galilée. Ce dernier a constaté que deux corps pouvant être assimilés à des points matériels de masses inertielles différentes acquièrent au cours d'une chute libre dans le vide une accélération identique. Le principe fondamental de la dynamique appliqué à un tel corps donne :

$$m_i \vec{a} = m_g \vec{g}$$

en notant $\vec{g}$ le champ de pesanteur¹. On a la même relation pour un autre corps dont on note les masses avec un prime. Le champ de pesanteur étant unique à l'endroit de la mesure, l'égalité des accélérations permet d'obtenir l'égalité suivante :

$$\frac{m_i}{m_g} = \frac{m'_i}{m'_g}$$

Il en résulte que masse inertielle et masse gravitationnelle sont proportionnelles. Le choix d'une même unité pour ces deux grandeurs conduit alors à leur égalité.

Ce résultat a été confirmé par la suite par différents physiciens qui ont amélioré la précision de cette proportionnalité. Ainsi en 1890, le hongrois Loránt von Eötvös a obtenu une précision de un milliardième à l'aide d'une adaptation de la balance de torsion imaginée par John Michell et utilisée par Henry Cavendish pour mesurer la constante de gravitation G en 1798. Dans les années 1960, Robert H. Dicke et

¹ On donnera une définition précise du champ de pesanteur dans un chapitre du cours de mécanique de la deuxième période consacré au mouvement d'une particule dans ce champ.

Vladimir B. Brazinsky ont obtenu une précision de 1 pour 10^{12} avec une méthode basée sur l'étude de l'accélération du Soleil.

Cette identité entre masse inertielle et masse gravitationnelle est connue en relativité comme le principe d'équivalence ; il s'agit d'une équivalence fondamentale.

e) Cas particulier d'un système à masse constante :

Dans ce cas, il est possible de « sortir » la masse de la dérivée de la quantité de mouvement puisqu'il s'agit d'une constante :

$$\frac{d\vec{p}_{/\mathcal{R}}(M)}{dt} = \frac{d(m\vec{v}_{/\mathcal{R}}(M))}{dt} = m \frac{d\vec{v}_{/\mathcal{R}}(M)}{dt} = \sum_i \vec{f}_i$$

Le principe fondamental de la dynamique s'écrit alors à l'aide du vecteur accélération sous la forme :

$$m\vec{a}_{/\mathcal{R}}(M) = \sum_i \vec{f}_i$$



Il faut se méfier de cette formulation qui ne s'applique qu'aux systèmes de masse constante. Le cas d'une fusée qui consomme du carburant et voit sa masse diminuer ou encore celui d'une goutte d'eau qui tombe dans une atmosphère contenant de la vapeur d'eau et grossit au cours du mouvement ne peuvent pas être traités avec cette relation.

Il est donc fortement conseillé de retenir l'expression générale du principe fondamental de la dynamique :

$$\frac{d\vec{p}_{/\mathcal{R}}(M)}{dt} = \sum_i \vec{f}_i$$

même si, dans le cadre de la mécanique du point, la formulation avec l'accélération est souvent suffisante.

3.3 Troisième loi de Newton : principe des actions réciproques

La troisième et dernière loi de Newton s'énonce de la manière suivante :

si un point matériel A exerce sur un point matériel B une force $\vec{f}_{A \rightarrow B}$, alors le point B exerce sur le point A une force $\vec{f}_{B \rightarrow A}$ telle que :

- les forces $\vec{f}_{A \rightarrow B}$ et $\vec{f}_{B \rightarrow A}$ s'exercent sur la même droite d'action à savoir la droite passant par A et B,
- $\vec{f}_{B \rightarrow A} = -\vec{f}_{A \rightarrow B}$.

Ce principe est parfois appelé « principe de l'action et de la réaction ».

Il est important de noter que ce principe ne concerne *a priori* que les points matériels. Dans ce cas, le problème est invariant par rotation autour de l'axe défini par les deux points A et B. Le principe de Curie² a pour conséquence que les forces sont portées par le même axe. On exclut la situation suivante :

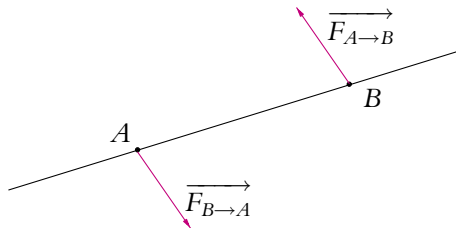


Figure 5.4 Couple de forces.

On verra lors de l'étude des dipôles électriques et magnétiques l'importance du caractère ponctuel des éléments du système pour que ce principe des actions réciproques soit vérifié.

Ces trois principes sont à la base de la dynamique classique ou newtonienne.

4. Résolution d'un problème de mécanique - Exemples

4.1 Méthode de résolution

Pour résoudre un problème de mécanique, il est judicieux de suivre les étapes suivantes dans l'ordre proposé par la suite. En effet, tous les points doivent apparaître et la succession qui va être énoncée satisfait à la logique du propos.

- La première chose à faire consiste à **définir le système**. Cela peut paraître inutile puisqu'on ne considère dans cette partie du cours qu'un point matériel mais cela deviendra fondamental dès lors qu'on envisagera des systèmes un peu plus complexes. En outre, il est toujours bon de préciser de quoi on parle.
- La deuxième étape consiste à **choisir le référentiel** qui sera utilisé. En première période, cette étape peut paraître superflue puisqu'on se limite à l'étude de mouvement dans des référentiels considérés comme galiléens. Cependant en seconde période, on verra qu'il est indispensable de connaître le caractère galiléen ou non du référentiel considéré avant d'établir le bilan des forces. On prendra donc dès maintenant la bonne habitude de préciser le référentiel dans lequel on effectue l'étude du mouvement.
- Il faut alors établir un **bilan des forces** précis et complet.

²Ce principe dit que les effets et les causes présentent les mêmes symétries. Il sera vu plus en détail dans la partie sur l'électromagnétisme.

- La dernière étape consiste à **choisir la méthode de résolution**. Pour l'instant, on ne dispose que d'une méthode : le principe fondamental de la dynamique mais on verra au chapitre suivant une méthode énergétique et en seconde période une troisième méthode basée sur un nouveau théorème. Le choix sera dicté par le problème et la facilité avec laquelle le problème se résout par l'une ou l'autre des approches. Cette étape est forcément la dernière puisqu'elle nécessite de connaître la réponse aux points précédents.

Toute résolution d'un problème mécanique devra donc suivre le schéma suivant dans cet ordre :

1. définition du système,
2. définition du référentiel,
3. bilan des forces appliquées,
4. choix de la méthode permettant la résolution.

4.2 Mouvement d'un point matériel dans le champ de pesanteur en l'absence de résistance de l'air

On étudie le mouvement d'un point matériel de masse m dans le champ de pesanteur $\vec{g}$.

a) Équation du mouvement

On va suivre la procédure de résolution d'un problème de mécanique décrite ci-dessus :

1. définition du système : le point matériel M ,
2. choix du référentiel : le référentiel terrestre que l'on considère galiléen,
3. bilan des forces : ne pas considérer la résistance de l'air revient à négliger tout phénomène de frottements. La seule force qui s'exerce dans ces conditions sur le point matériel est son poids : $\vec{P} = m\vec{g}$.
4. choix de la méthode de résolution : on utilise la seule méthode connue à ce stade à savoir le principe fondamental de la dynamique qui s'écrit :

$$m\vec{a} = \sum_i \vec{f}_i = m\vec{g}$$

soit

$$\vec{a} = \vec{g} \quad (5.1)$$

On peut constater que la masse n'intervient pas dans les équations du mouvement lorsqu'on néglige la résistance de l'air. On se souviendra que c'est cette observation expérimentale qui a permis d'établir l'égalité des masses inertielle et gravitationnelle.

Pour poursuivre l'analyse du mouvement, il est nécessaire de préciser les conditions initiales et par conséquent de distinguer le cas de la chute libre et celui du tir d'un point matériel.

b) Chute libre d'un point matériel

Les conditions initiales sont les suivantes : le point matériel est lâché d'une hauteur h sans vitesse initiale.

Équations horaires du mouvement

On remarquera que la seule force qui s'exerce est dirigée selon la verticale et que par conséquent les seules modifications du mouvement auront lieu le long de cet axe. On peut donc se limiter à l'étude le long de l'axe vertical.

On utilise une base de projection cartésienne en choisissant l'axe Oz vertical ascendant et on ne s'intéresse qu'au mouvement selon cet axe. L'équation du mouvement se traduit par :

$$\ddot{z} = -g$$

Soit, par intégrations successives, et en tenant compte des conditions initiales (vitesse nulle et altitude h) :

$$\dot{z} = -gt$$

et

$$z = -\frac{1}{2}gt^2 + h$$

On obtient donc un mouvement rectiligne uniformément accéléré le long de l'axe vertical Oz .

Caractéristiques du mouvement

On peut étudier différentes caractéristiques du mouvement :

- durée de chute T :

On l'obtient en résolvant l'équation $z = 0$ soit $-\frac{1}{2}gT^2 + h = 0$ d'où

$$T = \sqrt{\frac{2h}{g}}$$

- vitesse du point matériel quand il touche le sol : il s'agit de la valeur de la vitesse à l'instant $t = T$ soit

$$v = gT = \sqrt{2gh}$$

Cas d'une vitesse initiale verticale

On notera que lorsque la vitesse initiale est verticale, c'est-à-dire le long de l'axe selon lequel peuvent se produire des modifications du mouvement, le mouvement n'est pas intrinsèquement différent du cas de la chute libre : il est rectiligne et uniformément accéléré. En effet, on a alors :

$$\begin{aligned}\dot{z} &= -gt + v_0 \\ z &= -\frac{1}{2}gt^2 + v_0t + h\end{aligned}$$

La durée de chute est solution de l'équation

$$t^2 - 2\frac{v_0}{g}t - 2\frac{h}{g} = 0$$

Le discriminant réduit de cette équation est : $\Delta' = \left(\frac{v_0}{g}\right)^2 + 2\frac{h}{g}$. Il est positif, on a donc deux solutions réelles :

$$T = \frac{v_0}{g} \pm \sqrt{\left(\frac{v_0}{g}\right)^2 + 2\frac{h}{g}}$$

et seule la solution positive $T = \frac{v_0}{g} + \sqrt{\left(\frac{v_0}{g}\right)^2 + 2\frac{h}{g}}$ convient.

4.3 Trajectoire d'un tir dans le vide

On s'intéresse maintenant au cas où le point matériel est initialement animé d'une vitesse $\vec{v}_0$ faisant un angle α avec l'horizontale et où la position initiale est (x_0, z_0) .

Il suffit de limiter l'étude au plan xOz vertical dirigé par $\vec{g}$ et la vitesse initiale $\vec{v}_0$: les seules modifications du mouvement auront lieu dans ce plan puisqu'aucune force n'a d'action perpendiculairement à lui. Dans ce plan, les projections de l'équation (5.1) fournissent :

$$\begin{cases} \ddot{x} = 0 \\ \ddot{z} = -g \end{cases}$$

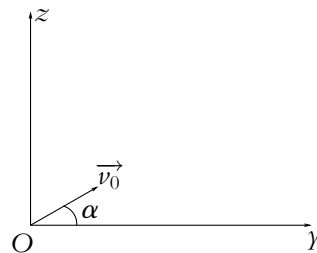


Figure 5.5 Conditions initiales d'un tir dans le vide.

Soit par intégrations successives et en tenant compte des conditions initiales :

$$\begin{cases} \dot{x} = v_0 \cos \alpha \\ \dot{z} = -gt + v_0 \sin \alpha \end{cases} \quad \text{puis} \quad \begin{cases} x = x_0 + v_0 t \cos \alpha \\ z = -\frac{1}{2}gt^2 + v_0 t \sin \alpha + z_0 \end{cases}$$

L'équation de la trajectoire s'obtient en éliminant le temps dans les équations horaires.

De l'expression de x en fonction du temps, on tire :

$$t = \frac{x - x_0}{v_0 \cos \alpha}$$

Puis en reportant dans l'expression de z :

$$z = -\frac{g}{2v_0^2 \cos^2 \alpha} x^2 + \left(\tan \alpha + \frac{gx_0}{v_0^2 \cos^2 \alpha} \right) x + z_0 - x_0 \tan \alpha - \frac{gx_0^2}{2v_0^2 \cos^2 \alpha}$$

La trajectoire est donc une parabole.

Par la suite, on prendra $x_0 = z_0 = 0$ par souci de simplification. L'équation de la trajectoire s'écrit alors :

$$z = -\frac{g}{2v_0^2 \cos^2 \alpha} x^2 + x \tan \alpha$$

► **Remarque :** On peut également intégrer l'équation vectorielle (5.1) et obtenir : $\vec{v} = \vec{g}t + \vec{v}_0$ puis $\vec{OM} = \frac{1}{2}\vec{g}t^2 + \vec{v}_0t + \vec{OM}_0$.

Cela permet de regrouper les deux cas et de ne projeter que pour étudier la trajectoire.

a) Hauteur maximale atteinte

Elle correspond au sommet de la parabole et s'obtient en résolvant l'équation :

$$\frac{dz}{dx} = \frac{-gx}{v_0^2 \cos^2 \alpha} + \tan \alpha = 0 \quad \text{soit} \quad x_m = \frac{\sin \alpha \cos \alpha}{g} v_0^2$$

L'altitude maximale est alors :

$$z_m = \frac{\sin^2 \alpha}{2g} v_0^2$$

b) Portée du tir

La portée du tir correspond à l'endroit où le point matériel retombe sur le sol qu'on suppose plan et horizontal (donc à l'altitude $z = 0$). On obtient l'expression de la portée en résolvant l'équation :

$$z = -\frac{g}{2v_0^2 \cos^2 \alpha} x^2 + x \tan \alpha = x \left(\tan \alpha - \frac{g}{2v_0^2 \cos^2 \alpha} x \right) = 0$$

On obtient deux solutions :

$$x = 0 \quad \text{et} \quad x = \frac{v_0^2 \sin 2\alpha}{g}$$

La solution $x = 0$ correspond à la position d'origine du tir qui appartient naturellement à la trajectoire. Il est normal de la trouver mais elle ne présente pas d'intérêt ici.

La portée du tir est donc :

$$x_p = \frac{v_0^2 \sin 2\alpha}{g}$$

On note qu'elle est maximale pour $\alpha = \frac{\pi}{4}$. C'est d'ailleurs l'angle avec lequel sautent les grenouilles pour retomber le plus loin possible.

c) Tir tendu ou tir en cloche

On va maintenant s'interroger sur la possibilité d'atteindre un point à la même altitude que celle où s'effectue le tir et situé à une distance donnée d de cette origine du tir. Le seul paramètre variable est l'angle de tir α .

D'après le paragraphe précédent, il faut chercher les valeurs de α vérifiant l'équation :

$$d = \frac{v_0^2 \sin 2\alpha}{g}$$

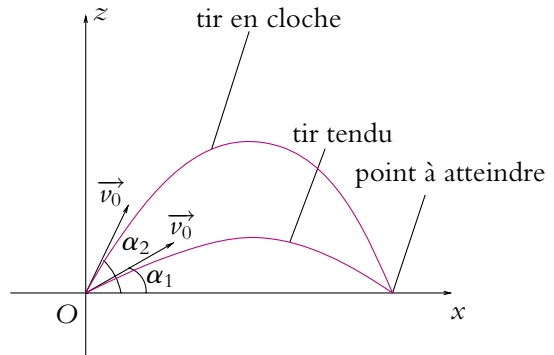


Figure 5.6 Les deux trajectoires possibles d'un tir dans le vide.

La résolution de cette équation conduit à deux solutions dans l'intervalle $[0, \frac{\pi}{2}]$:

$$\left\{ \begin{array}{l} \alpha = \frac{1}{2} \text{Arcsin} \frac{gd}{v_0^2} \\ \text{ou} \\ \alpha = \frac{\pi}{2} - \frac{1}{2} \text{Arcsin} \frac{gd}{v_0^2} \end{array} \right.$$

On obtient deux angles de tir possibles : la valeur la plus faible correspond au tir tendu tandis que la valeur la plus grande est celle du tir en cloche (voir figure 5.6).

d) Parabole de sûreté

On cherche à déterminer l'ensemble des points (x_0, z_0) qu'il est possible d'atteindre. Ce sera le cas si l'équation :

$$z_0 = -\frac{gx_0^2}{2v_0^2 \cos^2 \alpha} + x_0 \tan \alpha = -\frac{gx_0^2}{2v_0^2} (1 + \tan^2 \alpha) + x_0 \tan \alpha$$

admet des solutions³. On résout donc l'équation en $\tan \alpha$ suivante :

$$\frac{gx_0^2}{2v_0^2} \tan^2 \alpha - x_0 \tan \alpha + \frac{gx_0^2}{2v_0^2} + z_0 = 0$$

Le discriminant de cette équation du second degré en $\tan \alpha$ vaut

$$\Delta = x_0^2 - 4 \frac{gx_0^2}{2v_0^2} \left(\frac{gx_0^2}{2v_0^2} + z_0 \right)$$

L'équation n'admet des solutions réelles (seules acceptables ici) que si : $\Delta \geq 0$ donc si :

$$z_0 \leq \frac{v_0^2}{2g} - \frac{gx_0^2}{2v_0^2}$$

Un point pourra donc être atteint uniquement s'il se trouve dans le plan (xOz) sous la parabole d'équation :

$$z = \frac{v_0^2}{2g} - \frac{gx^2}{2v_0^2}$$

Cette parabole est appelée *parabole de sûreté*. Elle relie les sommets de toutes les trajectoires obtenues en faisant varier l'angle de lancement pour une vitesse initiale de module donné.

La figure suivante représente les trajectoires pour différents angles de tir avec une vitesse initiale de 10 m.s^{-1} ainsi que la parabole de sûreté en trait plus épais.

Pour un point situé sous la parabole de sûreté, on aura deux solutions pour l'angle α :

$$\tan \alpha = \frac{v_0^2 \pm \sqrt{v_0^4 - g^2 x_0 - 2gz_0 v_0^2}}{g}$$

³ On utilise le fait que $1 + \tan^2 \alpha = \frac{1}{\cos^2 \alpha}$: $1 + \frac{\sin^2 \alpha}{\cos^2 \alpha} = \frac{\cos^2 \alpha + \sin^2 \alpha}{\cos^2 \alpha} = \frac{1}{\cos^2 \alpha}$.

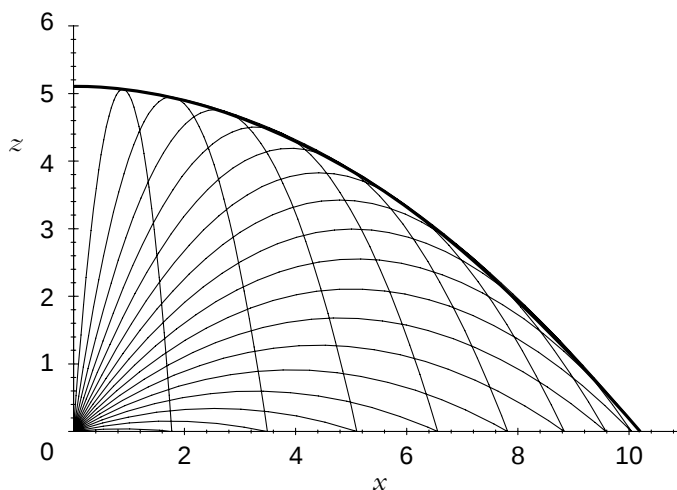


Figure 5.7 Parabole de sûreté.

Les angles de tir possibles pour atteindre le point considéré sont donc :

$$\alpha = \text{Arctan} \left(\frac{v_0^2 \pm \sqrt{v_0^4 - g_0^2 x_0 - 2gz_0 v_0^2}}{g} \right)$$

l'un correspondant au tir tendu, l'autre au tir en cloche.

On peut noter que lorsque le point considéré appartient à la parabole de sûreté, le discriminant est nul. Dans ce cas, il est possible de l'atteindre mais il n'y aura qu'une seule valeur possible pour l'angle de tir à savoir $\text{Arctan} \left(\frac{v_0^2}{g} \right)$.

4.4 Mouvement d'un point matériel dans le champ de pesanteur en présence de résistance de l'air

a) Influence de la résistance de l'air sur la chute d'un point matériel

On n'a pas tenu compte dans l'étude précédente de la résistance de l'air et plus généralement du milieu dans lequel s'effectue le mouvement. Le vide envisagé n'est qu'un cadre théorique simplificateur qu'on ne peut pas forcément conserver en pratique. Dans ce paragraphe, deux types de résistance de l'air seront envisagées : une force de frottement proportionnelle à la vitesse puis une force de frottement proportionnelle au carré de la vitesse.

Force de frottement proportionnelle à la vitesse

On ajoute au bilan des forces précédent la force de frottement $\vec{f} = -k\vec{v}$.

Le principe fondamental de la dynamique s'écrit alors :

$$m \vec{a} = -k \vec{v} + m \vec{g}$$

On suppose toujours que la vitesse initiale est nulle et que la masse est lâchée d'une hauteur h par rapport au sol. Toutes les forces sont verticales puisqu'initialement il n'y a que le poids et que ce dernier ne modifie la trajectoire que le long de la verticale. Le point matériel acquiert donc une vitesse verticale qui entraîne l'existence d'une force de frottement également verticale. Il suffit donc de considérer le mouvement en projection sur l'axe Oz vertical ascendant :

$$\frac{dv}{dt} + \frac{k}{m}v = -g$$

La solution de cette équation est la somme de la solution générale de l'équation homogène associée $v_H(t)$ et d'une solution particulière $v_P(t)$.

La résolution de l'équation homogène associée conduit à :

$$v_H(t) = C \exp\left(-\frac{k}{m}t\right)$$

où C est une constante. Dans la suite, on pose $\tau = \frac{m}{k}$, homogène à un temps.

La recherche d'une solution particulière est simple puisque le second membre est constant :

$$v_P(t) = -\frac{m}{k}g$$

La solution générale est donc :

$$v(t) = C \exp\left(-\frac{t}{\tau}\right) - \frac{m}{k}g$$

On détermine la constante à l'aide des conditions initiales à savoir

$$v(t=0) = 0 = C - \frac{m}{k}g$$

soit $C = \frac{m}{k}g$ et finalement on obtient :

$$v(t) = \frac{m}{k}g \left(\exp\left(-\frac{t}{\tau}\right) - 1 \right)$$

La position se déduit alors facilement par intégration par rapport au temps :

$$z = -\frac{m}{k}gt - \frac{m^2}{k^2}g \exp\left(-\frac{t}{\tau}\right) + K$$

et la constante K s'obtient en tenant compte de la position initiale :

$$z(t=0) = h = -\frac{m^2}{k^2}g + K \quad \text{soit} \quad K = h + \frac{m^2}{k^2}g$$

et finalement

$$z(t) = -\frac{m}{k}gt + \frac{m^2}{k^2}g \left(1 - \exp\left(-\frac{t}{\tau}\right)\right) + h$$

Vitesse limite

On peut à partir des résultats précédents établir l'existence d'une vitesse limite obtenue quand le temps tend vers l'infini :

$$v_l = -\frac{m}{k}g$$

On peut noter qu'il s'agit de la valeur obtenue quand l'accélération s'annule : quand la vitesse tend vers une constante, l'accélération est nulle. Elle est indépendante des conditions initiales. Cependant elle n'est pas obtenue instantanément mais après un régime transitoire dont la durée caractéristique est la constante τ . La constante τ s'interprète donc physiquement comme le temps caractéristique d'établissement du régime pour lequel la vitesse est égale à la vitesse limite.

Force de frottement proportionnelle au carré de la vitesse

On considère maintenant que la force de frottement se met sous la forme

$$-kv \vec{v}$$

et le principe fondamental de la dynamique devient :

$$m \frac{d\vec{v}}{dt} = m\vec{g} - kv \vec{v}$$

On considère toujours le cas particulier où la vitesse initiale est nulle. Le problème est alors unidimensionnel pour les mêmes raisons qu'avec une force de frottement proportionnelle à la vitesse. On a donc simplement l'équation différentielle scalaire suivante à résoudre :

$$\frac{dv}{dt} + \frac{k}{m}v^2 = g$$

Cette équation n'est pas linéaire, on ne peut donc pas la résoudre comme la précédente.

Vitesse limite

Comme dans le cas d'un frottement proportionnel à la vitesse, le point matériel atteint une vitesse limite vérifiant :

$$\frac{k}{m}v_l^2 = g$$

car l'accélération tend vers 0 lorsque la vitesse tend vers une constante. La vitesse limite a donc pour expression :

$$v_l = \sqrt{\frac{mg}{k}}$$

On note qu'elle est indépendante des conditions initiales.

Expression de la vitesse

On peut alors introduire la vitesse limite dans l'équation différentielle qui devient :

$$\frac{dv}{dt} = \frac{k}{m} (v_l^2 - v^2)$$

C'est une équation différentielle à variables séparables :

$$\frac{dv}{v_l^2 - v^2} = \frac{k}{m} dt$$

Une solution classique⁴ pour obtenir la solution de cette équation différentielle consiste alors à décomposer $\frac{1}{v_l^2 - v^2}$ en éléments simples :

$$\frac{1}{v_l^2 - v^2} = \frac{1}{2v_l(v_l - v)} + \frac{1}{2v_l(v_l + v)}$$

et on peut réécrire l'équation différentielle sous la forme :

$$\frac{dv}{v_l - v} + \frac{dv}{v_l + v} = \frac{2kv_l}{m} dt = 2\frac{t}{\tau}$$

en notant $\tau = \frac{m}{kv_l} = \sqrt{\frac{m}{kg}}$. On obtient par intégration entre $t = 0$ et t et en tenant compte des conditions initiales :

$$-\ln \left| \frac{v_l - v}{v_l} \right| + \ln \left| \frac{v_l + v}{v_l} \right| = \ln \left| \frac{v_l + v}{v_l - v} \right| = 2\frac{t}{\tau} \quad \text{soit} \quad \frac{v_l + v}{v_l - v} = \exp \left(\frac{2t}{\tau} \right)$$

On en déduit :

$$v = v_l \operatorname{th} \frac{t}{\tau} = \sqrt{\frac{mg}{k}} \operatorname{th} \frac{t}{\tau}$$

⁴ L'autre façon de procéder consiste à connaître une primitive d'expression du type $\frac{1}{a^2 - x^2}$.

La constante τ s'interprète donc comme la durée caractéristique au bout de laquelle le point matériel atteint sa vitesse limite.

La position s'obtient en intégrant cette expression en fonction du temps :

$$z = h + \int_0^t v(u) du = h + \int_0^t \sqrt{\frac{mg}{k}} \operatorname{th} \frac{u}{\tau} du = h + \frac{m}{k} \ln \left| \operatorname{ch} \frac{t}{\tau} \right|$$

On peut également obtenir facilement l'expression de z en fonction de la vitesse en remarquant que :

$$\frac{dv}{dt} = \frac{dv}{dz} \frac{dz}{dt} = \frac{dv}{dz} v = \frac{1}{2} \frac{d(v^2)}{dz}$$

En reportant cette expression dans l'équation du principe fondamental de la dynamique, on obtient :

$$\frac{d(v^2)}{dz} = 2 \frac{k}{m} (v_l^2 - v^2)$$

En posant $u = v^2$, on obtient l'équation différentielle **linéaire** du premier ordre en u :

$$\frac{du}{dz} + \frac{2k}{m} u = \frac{2kv_l^2}{m}$$

dont la solution est :

$$u = v^2 = v_l^2 + A \exp\left(-\frac{z}{H}\right)$$

en posant $H = \frac{m}{2k}$ qui correspond à une distance caractéristique. Or à $t = 0$, $v = 0$

et $z = h$ donc $0 = v_l^2 + A \exp\left(-\frac{h}{H}\right)$ soit $A = -v_l^2 \exp\left(\frac{h}{H}\right)$ et

$$v^2 = v_l^2 \left(1 - \exp\left(-\frac{z-h}{H}\right)\right)$$

soit en inversant :

$$z = h - H \ln \frac{v_l^2 - v^2}{v_l^2}$$

La distance H s'interprète donc comme la distance caractéristique pour laquelle le point matériel atteint la vitesse limite.

Comparaison et discussion

L'objet de ce paragraphe est de comparer les deux types de frottement envisagés.

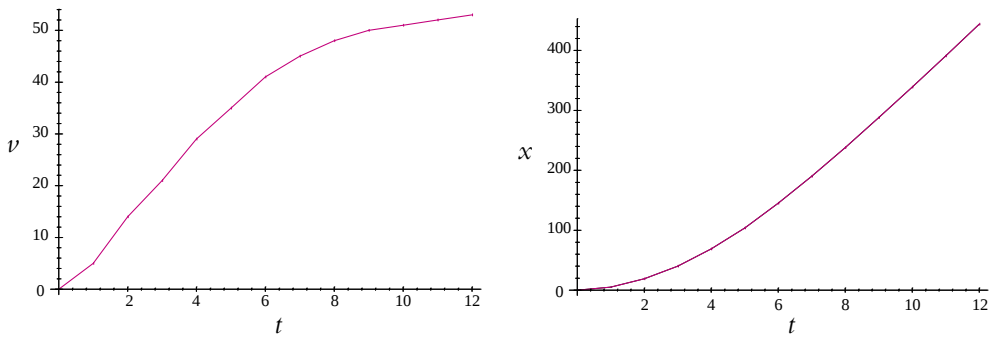


Figure 5.8 Vitesse et position d'une chute en parachute.

- Dans les deux cas, il existe une vitesse limite dont la valeur n'est pas la même :

$$\begin{cases} v_l = \frac{mg}{k} & \text{pour un frottement proportionnel à la vitesse,} \\ v_l = \sqrt{\frac{mg}{k}} & \text{pour un frottement proportionnel au carré de la vitesse.} \end{cases}$$

Les figures ci-dessus montrent l'évolution temporelle de la vitesse (en m.s^{-1}) et de la distance parcourue (en m) lors du saut d'un parachutiste. On observe l'existence de la vitesse limite et l'adéquation de ce qui vient d'être établi avec l'expérience.

- Un frottement proportionnel à la vitesse convient bien lorsque les vitesses sont faibles : il est toujours possible dans ce cas de linéariser. En revanche, lorsque les vitesses deviennent plus importantes, il est nécessaire d'envisager un frottement proportionnel à une puissance de la vitesse. On a traité ici le cas du carré mais d'autres valeurs de puissance pourraient être considérées.
- La méthode employée dans le cas d'un frottement proportionnel au carré de la vitesse est intéressante : l'introduction d'une valeur limite pour la vitesse permet de « ramener » une équation différentielle *a priori* complexe à résoudre à une équation « classique ». Plus que le résultat, il peut être utile de retenir l'idée d'introduire un paramètre limite pour simplifier la résolution.

Des précisions sur l'expression de la force de frottement seront apportées dans le cours de mécanique des fluides de seconde année PC et PSI.

b) Influence de la résistance de l'air sur le tir d'un point matériel

On se limite au cas d'une force de frottement proportionnelle à la vitesse pour modéliser la résistance de l'air. L'équation du mouvement est la même que celle de la chute libre :

$$m \vec{a} = -k \vec{v} + m \vec{g} \quad \text{soit} \quad \frac{d\vec{v}}{dt} + \frac{k}{m} \vec{v} = \vec{g}$$

On peut résoudre cette équation différentielle du premier ordre soit scalairement après projection dans une base (la base cartésienne est bien adaptée au problème) soit vectoriellement, ce qui va être fait ici.

La solution s'écrit comme la somme de la solution générale de l'équation homogène associée $\vec{v}_H$ et d'une solution particulière $\vec{v}_P$.

La résolution de l'équation homogène associée conduit à :

$$\vec{v}_H = \vec{C} \exp\left(-\frac{t}{\tau}\right)$$

en notant $\vec{C}$ une constante vectorielle et en posant $\tau = \frac{m}{k}$.

La recherche d'une solution particulière est simple puisque le second membre est une constante :

$$\vec{v}_P = \frac{m}{k} \vec{g}$$

La solution globale est donc :

$$\vec{v} = \vec{C} \exp\left(-\frac{t}{\tau}\right) + \frac{m}{k} \vec{g}$$

On détermine la constante à l'aide des conditions initiales, à savoir

$$\vec{v}(t=0) = \vec{v}_0 = \vec{C} + \frac{m}{k} \vec{g}$$

soit $\vec{C} = \vec{v}_0 - \frac{m}{k} \vec{g}$ et finalement on obtient :

$$\vec{v} = \vec{v}_0 \exp\left(-\frac{t}{\tau}\right) + \frac{m}{k} \vec{g} \left(\exp\left(-\frac{t}{\tau}\right) - 1\right)$$

La position se déduit alors facilement par intégration par rapport au temps :

$$\vec{OM} = -\frac{m}{k} \vec{v}_0 \exp\left(-\frac{t}{\tau}\right) + \frac{m}{k} t \vec{g} + \frac{m^2}{k^2} \vec{g} \exp\left(-\frac{t}{\tau}\right) + \vec{K}$$

et la constante $\vec{K}$ s'obtient en tenant compte de la position initiale :

$$\vec{OM}(t=0) = \vec{0}$$

et finalement

$$\vec{OM} = \frac{m}{k} \vec{g} t + \frac{m}{k} \left(\frac{m}{k} \vec{g} - \vec{v}_0\right) \left(\exp\left(-\frac{t}{\tau}\right) - 1\right)$$

Vitesse limite

On peut à partir des résultats précédents établir l'existence d'une vitesse limite obtenue quand le temps tend vers l'infini :

$$\vec{v}_l = \frac{m}{k} \vec{g}$$

On peut noter qu'il s'agit de la valeur obtenue quand l'accélération s'annule. Elle est indépendante des conditions initiales et est atteinte au bout de quelques τ .

Trajectoire

On laisse au lecteur le soin d'établir que l'équation de la trajectoire est :

$$z = -\frac{m^2 g}{k^2} \ln \left(\frac{mv_0 \cos \alpha}{mv_0 \cos \alpha - kx} \right) + \left(\frac{mg}{k} + v_0 \sin \alpha \right) \frac{x}{v_0 \cos \alpha}$$

On peut noter que lorsque le temps t tend vers l'infini alors x tend vers $\frac{mv_0 \cos \alpha}{k}$, ce qui établit l'existence d'une asymptote verticale au mouvement.

On peut également établir que la trajectoire atteint une altitude maximale pour

$$x = \frac{mv_0^2 \cos^2 \alpha}{mg + kv_0 \cos \alpha}$$

L'allure de la trajectoire est donnée ci-contre.

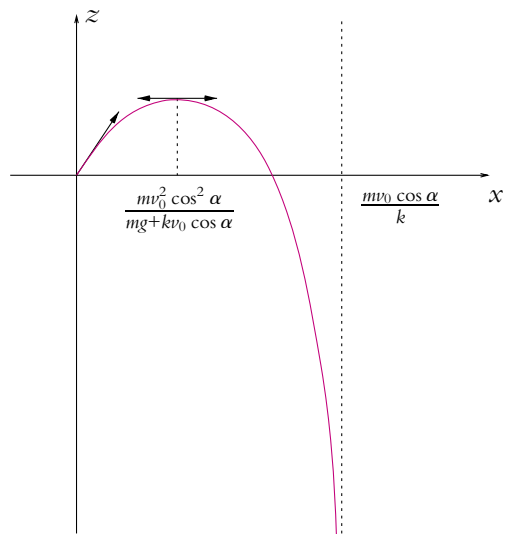


Figure 5.9 Trajectoire d'un tir avec un frottement proportionnel à la vitesse.

4.5 Mouvement d'une masse attachée à un ressort dont l'autre extrémité est fixe

On s'intéresse au mouvement d'un point matériel M de masse m attaché à un ressort parfaitement élastique dont l'autre extrémité est fixe. On note k la raideur du ressort et l_0 sa longueur à vide.

On distingue les deux cas d'un mouvement horizontal et d'un mouvement vertical.

a) Cas d'un mouvement horizontal

Dans ce paragraphe, on suppose que la masse m coulisse sans frottement le long de l'axe horizontal Ox .

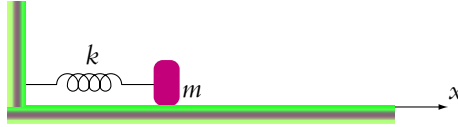


Figure 5.10 Pendule élastique horizontal.

On applique les quatre étapes de résolution d'un problème mécanique :

1. Système : point matériel de masse m .
2. Référentiel : celui du laboratoire considéré comme galiléen.
3. Bilan des forces :
 - poids $m\vec{g}$,
 - réaction $\vec{R}$ du support horizontal le long duquel la masse se déplace, cette force est perpendiculaire au support du fait de l'absence de frottement,
 - force de rappel du ressort $\vec{f} = -k(l - l_0)\vec{u}_x$ où l désigne la longueur du ressort à l'instant considéré.
4. Choix de la méthode de résolution : principe fondamental de la dynamique.

On a donc :

$$m\vec{a} = m\vec{g} + \vec{R} - k(l - l_0)\vec{u}_x$$

La projection sur la verticale donne : $R = mg$ qui traduit l'équilibre de la masse suivant cette direction et le fait que le mouvement n'est qu'horizontal.

À l'équilibre, la longueur du ressort est sa longueur à vide. On choisit sur l'axe Ox l'origine au niveau de la position de la masse à l'équilibre, c'est-à-dire ici quand le ressort est à vide. La longueur du ressort s'écrit donc : $l = l_{eq} + x = l_0 + x$: $x = l - l_{eq}$ est le déplacement de la masse m par rapport à sa position d'équilibre (cela correspond ici à l'allongement du ressort car $x = l - l_0$). La projection du principe fondamental de la dynamique sur l'axe Ox fournit : $m\ddot{x} = -kx$ soit :

$$\ddot{x} + \frac{k}{m}x = 0$$

b) Cas d'un mouvement vertical

On suppose maintenant que le mouvement est vertical.

1. Système : point matériel de masse m .
2. Référentiel : celui du laboratoire considéré comme galiléen.

3. Bilan des forces :

- poids $m\vec{g}$,
- force de rappel du ressort $\vec{f} = -k(l - l_0)\vec{u}_x$ où l désigne la longueur du ressort à l'instant considéré.

4. Choix de la méthode de résolution : principe fondamental de la dynamique.

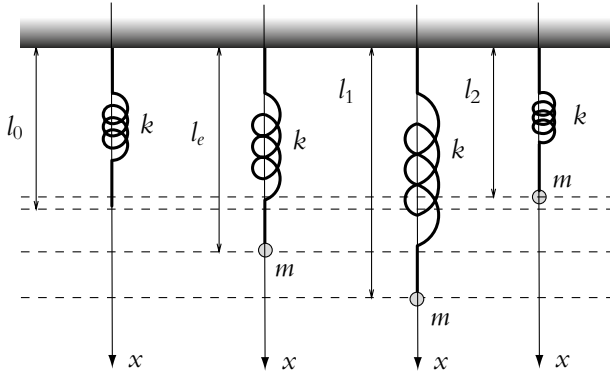


Figure 5.11 Masse suspendue à un ressort.

$$m\vec{a} = m\vec{g} - k(l - l_0)\vec{u}_x$$

À l'équilibre, on obtient en projection sur l'axe vertical :

$$-k(l_e - l_0) + mg = 0$$

Le ressort a donc une longueur à l'équilibre égale à

$$l_e = l_0 + \frac{mg}{k}$$

On choisit toujours l'origine de l'axe vertical Ox à la position d'équilibre à savoir quand le ressort a une longueur l_e . L'abscisse x est alors le déplacement de la masse m par rapport à sa position d'équilibre :

$$l = l_e + x$$

Il s'agit d'une quantité algébrique pouvant être négative ou positive (les deux cas sont représentés sur la figure ci-dessus). La projection du principe fondamental de la dynamique le long de cet axe conduit à la relation :

$$m\ddot{x} = mg - k(l - l_0) = mg - k(l_e + x - l_0)$$

En reportant l'expression obtenue pour l_e , on obtient l'équation différentielle suivante :

$$\ddot{x} + \frac{k}{m}x = 0$$



Il est important de noter la méthode employée ici qui consiste à introduire la position d'équilibre aussi bien pour le mouvement horizontal que pour le mouvement vertical et à définir x non pas comme l'allongement du ressort mais comme l'écart par rapport à la position d'équilibre. L'intérêt de cette approche est d'éliminer toutes les constantes de l'équation du mouvement et d'obtenir dans les deux cas la même équation :

$$\ddot{x} + \frac{k}{m}x = 0$$

Il s'agit de l'équation d'un oscillateur harmonique dont la solution s'écrit :

$$x(t) = A \cos \omega_0 t + B \sin \omega_0 t$$

en notant $\omega_0 = \sqrt{\frac{k}{m}}$ et A et B deux constantes. Si on suppose qu'initialement on a $x(0) = x_0$ et $\dot{x}(0) = v_0$, on en déduit :

$$x(t) = x_0 \cos \omega_0 t + \frac{v_0}{\omega_0} \sin \omega_0 t$$

On développera l'étude de ce mouvement dans un chapitre ultérieur consacré aux oscillateurs mécaniques.

A. Applications directes du cours

1. Masse en rotation

On considère un fil sans masse de longueur l , dont l'une des extrémités est reliée à l'axe vertical Oz au point A et l'autre à un point M de masse m (figure 5.12). L'axe Oz est en rotation à la vitesse angulaire constante ω .

Déterminer l'angle α , la vitesse de M et la tension du fil en fonction de ω , m et g l'accélération de la pesanteur.

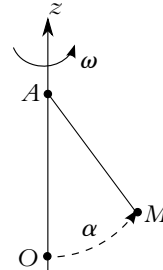


Figure 5.12

2. Bond sur la Lune

Dans l'album de Tintin, « On a marché sur la Lune », le Capitaine Haddock est étonné de pouvoir faire un bond beaucoup plus grand que sur Terre : c'est le sujet de cet exercice.

Le personnage saute depuis le sol lunaire avec une vitesse initiale v_0 faisant un angle $\alpha = 30^\circ$ avec le sol. On note g_l l'accélération de la pesanteur à la surface de la lune. En l'absence d'atmosphère, on peut considérer qu'il n'y a aucune force de frottement.

- Déterminer l'expression de la distance parcourue au cours du saut en fonction de v_0 , α et g_l .
- La gravité sur la Lune est environ six fois moins forte que sur la Terre. Quelle sera la distance parcourue par le sauteur sur la Lune, si elle est de $d = 1,50$ m sur la Terre ?

3. Masse attachée à deux ressorts

On considère un point M de masse m attaché à deux ressorts identiques verticaux, de constante de raideur k et de longueur à vide ℓ_0 . (figure 5.13). les deux autres extrémités O et O' des ressorts sont fixes et espacées d'une distance L . On définit l'axe Oz vertical ascendant.

- Déterminer la position d'équilibre z_e de M .
- Déterminer l'équation différentielle à laquelle satisfait $z(t)$. On écrira cette équation en fonction de $\omega_0 = \sqrt{2k/m}$ et z_e .
- On écarte M d'une hauteur a par rapport à sa position d'équilibre, et on le lâche sans vitesse. Déterminer $z(t)$.

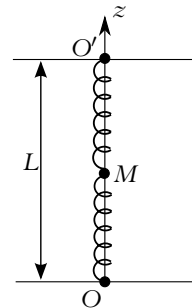


Figure 5.13

4. Charge soulevée par une grue

Une grue de chantier, de hauteur h doit déplacer d'un point à un autre du chantier une charge M de masse m supposée ponctuelle.

On appelle A le point d'attache du câble sur le chariot de la grue.

- Le point A est à la verticale de M posée sur le sol. Déterminer la tension du câble lorsque M décolle.

2. L'enrouleur de câble de la grue remonte le câble avec une accélération a_v constante. Déterminer la tension du câble. Conclusion.

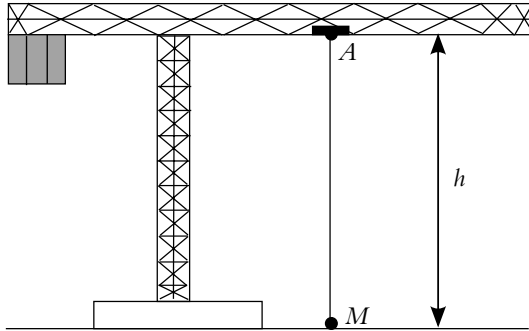


Figure 5.14

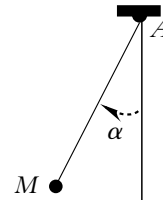


Figure 5.15

3. La montée de M est stoppée à mi-hauteur mais le chariot A se met en mouvement vers la droite (figure 5.14) avec une accélération a_h constante.

- Quelle est l'accélération de M sachant que M est alors immobile par rapport à A ?
- Déterminer l'angle que fait le câble avec la verticale en fonction de m , g , a_h ainsi que la tension du câble.

B. Exercices et problèmes

1. Quelques tirs de basket-ball

On étudie les tirs de basket-ball de manière simplifiée. On suppose que le joueur est face au panneau à une distance D de ce dernier. Le cercle du panier est situé à une hauteur $H = 3,05$ m au-dessus du sol et on assimilera dans un premier temps le cercle à un point situé sur le panneau. De même, le ballon sera considéré comme ponctuel. On néglige les frottements fluides de l'air. Le joueur tire d'une hauteur $h = 2,00$ m au-dessus du sol en imposant une vitesse initiale $\vec{v}_0$ faisant un angle α avec l'horizontale.

- Établir l'équation du mouvement du ballon lors du tir.
- En déduire les équations horaires de ce mouvement.
- Déterminer l'équation de la trajectoire du ballon.
- On suppose que le module de la vitesse initiale est fixé. Donner l'équation à vérifier par l'angle α pour que le panier soit marqué. On la mettra sous la forme d'une équation du second degré en $\tan \alpha$.
- Montrer que cette équation n'admet des solutions que si le module v_0 de la vitesse initiale vérifie une inéquation du second degré en v_0^2 .
- En déduire l'existence d'une valeur minimale de v_0 pour que le panier soit marqué.

7. Faire l'application numérique pour un lancer franc (la distance D vaut alors 4,60 m) puis pour un panier à trois points (la distance D vaut alors 6,25 m selon les règles de la Fédération Internationale de Basket-Ball).
8. Si la condition précédente est vérifiée, donner l'expression de $\tan \alpha$ et en déduire qu'il existe deux angles possibles pour marquer le panier.
9. Donner les valeurs numériques des angles α permettant de marquer un lancer franc en supposant que $v_0 = 10,0 \text{ m.s}^{-1}$.
10. Dans la suite, on suppose que l'angle de tir est fixé. Déterminer l'expression de la vitesse initiale v_0 à imposer pour marquer le panier.
11. Faire l'application numérique pour un lancer franc et un angle de tir $\alpha = 70^\circ$.
12. On s'intéresse maintenant à l'influence de la dimension du ballon et du cercle du panier. On rappelle que le rayon du ballon est de 17,8 cm et que le diamètre du cercle mesure 45,0 cm. On note x la distance du centre du ballon au panneau à hauteur du cercle du panier. Donner l'intervalle des valeurs de x permettant de marquer le panier.
13. Quelle équation doit vérifier x ? On pourra la laisser sous la forme d'une équation du second degré en $(D + x)$.
14. En déduire la condition que doivent vérifier v_0 , α , g , H et h pour que cette équation ait des solutions.
15. En déduire les valeurs d'angle de tir possibles en supposant la vitesse initiale v_0 fixée. On fera l'application numérique pour $v_0 = 10,0 \text{ m.s}^{-1}$.
16. Établir l'existence d'une valeur minimale de v_0 à angle de tir fixé. On fera l'application numérique pour $\alpha = 70^\circ$.
17. Cette condition étant vérifiée, donner l'expression de x .

2. Chute avec frottement quadratique, d'après ESIM 2001

Le repère terrestre est pris galiléen et le module de l'accélération de la pesanteur est $g = 9,81 \text{ m.s}^{-2}$.

L'axe Oz est orienté vers le bas et seul est étudié le mouvement de chute d'un corps de masse m . La trajectoire est donc rectiligne. Dans l'air, ce corps est alors soumis, selon Oz , à une force de frottement visqueux de norme $k\nu^2$ avec $k = \frac{1}{2}C_x\rho_0 S$ où $\rho_0 = 1,22 \text{ kg.m}^{-3}$ est la masse volumique de l'air, ν le module de la vitesse du corps, S est une section et C_x un coefficient aérodynamique, tous deux caractéristiques de la forme du corps.

Le poids et le frottement sont les deux seules forces prises en compte.

1. Comment peut-on comprendre qualitativement l'existence d'une vitesse limite ? Exprimer cette vitesse limite v_l en fonction des données de l'énoncé.
2. a) Écrire l'équation différentielle donnant la vitesse ν du corps en y faisant figurer comme seules constantes v_l et g . Quelle est la caractéristique de cette équation ?
b) L'intégrer directement après avoir séparé les variables, et exprimer la vitesse $\nu(t)$ à l'aide d'une fonction hyperbolique et des constantes v_l et $\tau = v_l/g$ en choisissant comme instant initial $t = 0$. Quelles sont la dimension et la signification physique de τ ?

On rappelle que la dérivée de la fonction $f(x) = \tanh^{-1}\left(\frac{x}{a}\right)$ est $f'(x) = \frac{1}{a} \frac{1}{1 - (x/a)^2}$.

c) Sachant que le corps est lâché à l'instant initial de la cote $z = 0$, donner l'expression de sa cote ultérieure $z(t)$ en fonction des constantes v_l et τ . Ce résultat dépend-il de la masse m ? S'il n'y avait pas de frottement, est ce que $z(t)$ dépendrait de la masse m ?

3. Application numérique : elle concerne des corps sphériques de diamètre d pour lesquels $C_x \simeq 1$ et $S = \pi d^2/4$ S étant la plus grande section droite du corps perpendiculaire à la direction du mouvement. Elle est menée parallèlement pour :

- une balle de tennis (objet 1) pour laquelle $m_1 = 57,6$ g et $d_1 = 6,70$ cm ;
- une boule de pétanque (objet 2) pour laquelle $m_2 = 700$ g et $d_2 = 7,60$ cm.

a) Calculer τ_1 et τ_2 d'une part, et $v_{l,1}$ et $v_{l,2}$ de l'autre. À quel paramètre est essentiellement due cette différence ?

Lâchées simultanément de la même hauteur, quelle est, de la balle de tennis et de la boule de pétanque, celle qui touche le sol en premier ?

b) La hauteur de chute est fixée à $h = 1,80$ m au dessus du sol (hauteur d'homme).

- Calculer les deux temps de chute t_1 et t_2 puis exprimer en millisecondes la différence $\Delta t = t_1 - t_2$.
- Avec combien de centimètres d'avance le premier corps arrive-t-il avant le second ? Ces différences sont-elles suffisantes pour être notables ?

c) La hauteur de chute est fixée à $h = 10$ m au dessus du sol (4ème étage d'un immeuble). Reprendre les applications numériques de la question précédente et conclure.

3. Toboggan

Une personne participant à un jeu télévisé doit laisser glisser un paquet sur un toboggan, à partir du point A de manière à ce qu'il soit réceptionné, au point B par un chariot se déplaçant le long de l'axe Ox . On néglige les frottements de l'air. Le toboggan exerce sur le paquet non seulement une réaction normale $\vec{R}_N$ mais aussi une réaction tangentielle $\vec{R}_t$ (frottement solide) telle que $\|\vec{R}_t\| = f\|\vec{R}_N\|$ avec $f = 0,5$.

Le chariot part du point C_0 à l'instant $t = 0$, vers la gauche (figure 5.16) avec une vitesse $v_c = 0,5 \text{ m.s}^{-1}$. Arrivé en C_B , il s'arrête pendant un intervalle de temps $\delta t = 1$ s.

La hauteur du toboggan est $h = 4$ m, la distance d est égale à h et $C_0C_B = 2$ m. La masse du paquet est $m = 10$ kg et $g = 9,8 \text{ m.s}^{-2}$.

1. On pose $X = AP$ où P est la position du paquet à l'instant t . Déterminer $X(t)$. En déduire le temps mis par le paquet pour arriver en B .

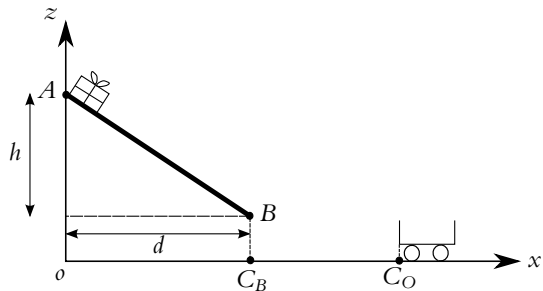


Figure 5.16

2. À quel instant le joueur doit-il lâcher le paquet, sans vitesse initiale, pour qu'il tombe dans le chariot ?

4. Élastique de saut

Un fabricant d'élastiques de saut donne les caractéristiques suivantes pour un type d'élastique :

Type **M** : pour un poids de sauteur de 65 à 95 kg.

- Longueur à vide (notée ℓ_0) : de 6 m à 50 m pour des sites de saut de 30 m à 250 m.
- Tension appliquée sur un élastique pour un allongement de 100 % : 200 kg.
- Tension appliquée sur un élastique pour un allongement de 200 % : 325 kg.

On prendra $g = 9,8 \text{ m.s}^{-2}$.

1. Traduire la fiche technique en langage correct en physique.

2. a) La tension de l'élastique obéit-elle à une loi type ressort : $\overline{T} = k(\ell - \ell_0)$ où T est la tension appliquée et ℓ la longueur de l'élastique ?

b) On supposera dans la suite que la tension est de la forme $\overline{T} = k(\ell - \ell_0)$. À partir des données, déterminer une valeur moyenne de la constante de raideur de l'élastique k pour un élastique de $\ell_0 = 6 \text{ m}$,

3. On veut savoir à quelle hauteur remonterait une masse test M de masse minimale (ici $m = 65 \text{ kg}$) si elle est lâchée sans vitesse initiale, l'élastique tendu au maximum. On prendra l'exemple d'un élastique de $\ell_0 = 6 \text{ m}$, dont le point d'attache est à la hauteur $h = 18 \text{ m}$ au-dessus du point de départ.

a) Déterminer l'expression de l'altitude $z(t)$ et de la vitesse $\dot{z}(t)$ tant que l'élastique est tendu.

b) Calculer l'instant t_1 pour lequel l'élastique n'est plus tendu. En déduire la vitesse de M en cet instant.

c) On ne tient pas compte des frottements de l'air. Déterminer la hauteur maximale atteinte par l'objet. Conclusion.

4. On réalise maintenant un saut normal, à partir du point d'attache, sans vitesse initiale avec les mêmes caractéristiques qu'à la question précédente.

a) Déterminer le point le plus bas atteint par le sauteur et l'accélération ressentie en ce point. Conclusion.

b) Après plusieurs oscillations, le sauteur s'immobilise. À quelle hauteur ?

5. Jeux aquatiques

Un baigneur (masse $m = 80 \text{ kg}$) saute d'un plongoir situé à une hauteur $h = 10 \text{ m}$ au-dessus de la surface de l'eau. On considère qu'il se laisse chuter sans vitesse initiale et qu'il est uniquement soumis à la force de pesanteur (on prendra $g = 10 \text{ m.s}^{-2}$) durant la chute. On note Oz , l'axe vertical descendant, O étant le point de saut.

1. Déterminer la vitesse v_e d'entrée dans l'eau ainsi que le temps de chute t_e . Application numérique.

2. Lorsqu'il est dans l'eau, le baigneur ne fait aucun mouvement. Il subit en plus de la pesanteur :

- une force de frottement $\vec{f}_f = -k\vec{v}$ ($\vec{v}$ étant la vitesse et $k = 250 \text{ kg.s}^{-1}$);
 - la poussée d'Archimède $\vec{\Pi} = -\frac{m}{d_h}\vec{g}$ ($d_h = 0,9$ est la densité du corps humain).
- a) Établir l'équation différentielle à laquelle obéit la vitesse en projection sur Oz , notée v_z . On posera $\tau = m/k$.
 - b) Intégrer cette équation en prenant comme nouvelle origine des temps $t = t_c$.
 - c) Déterminer la vitesse limite v_L (<0) en fonction de m , k , g et d_h . Application numérique.
 - d) Exprimer la vitesse v_z en fonction de v_c , $|v_L|$ et t . Déterminer à quel instant t_1 le baigneur commence à remonter.
 - e) En prenant la surface de l'eau comme nouvelle origine de l'axe Oz , exprimer $z(t)$. En déduire la profondeur maximale pouvant être atteinte.
 - f) En fait, il suffit que le baigneur arrive au fond de la piscine avec une vitesse de l'ordre de 1 m.s^{-1} pour pouvoir se repousser avec ses pieds : à quel instant t_2 atteint-il cette vitesse et quelle est la profondeur minimale du bassin?
3. Le même baigneur décide maintenant d'effectuer un plongeon. On suppose qu'il entre dans l'eau avec un angle $\alpha = 60^\circ$ par rapport à l'horizontale et une vitesse $v_0 = 8 \text{ m.s}^{-1}$. Les forces qui s'exercent sur lui sont les mêmes que précédemment mais le coefficient k est divisé par deux en raison d'une meilleure pénétration dans l'eau.
- On repère le mouvement par les axes Ox (axe horizontal de même sens que $\vec{v}_0$) et Oz (vertical descendant comme précédemment); le point O est le point de pénétration dans l'eau.
- a) Déterminer les projections des équations du mouvement sur Ox et Oz .
 - b) En déduire les composantes de la vitesse dans l'eau en fonction du temps. Existe-t-il une vitesse limite?
 - c) Le plongeur peut-il atteindre le fond de la piscine situé à 4 m ?

6. Décollage d'un véhicule

Une voiture, assimilé à un point matériel M de masse $m = 1000 \text{ kg}$ amorce une descente en A à la vitesse $v_0 = 125 \text{ km.h}^{-1}$ (figure 5.17). On assimile le début de la descente de A à B à un arc de cercle, de centre O , de rayon $R = 130 \text{ m}$ et d'angle $\alpha = 15^\circ$.

On suppose que durant la descente la force motrice de la voiture est tangente à la route et de valeur algébrique $\bar{F}$ constante (positive pour une accélération et négative pour un freinage). On néglige les frottements sur la route.

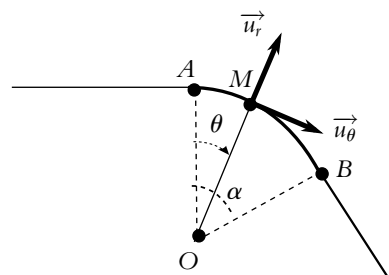


Figure 5.17

1. Déterminer les équations du mouvement de $M(R, \theta)$ en projection sur la base polaire de centre O : $(\vec{u}_r, \vec{u}_\theta)$.
2. Après avoir multiplié l'équation sur $\vec{u}_\theta$ par $\dot{\theta}$, l'intégrer.

3. Déterminer la réaction normale R_n en fonction de R , θ , g (pesanteur), m , v_0 et $\overline{F}$.
4.
 - a) Déterminer l'expression vérifiée par l'angle θ_d pour lequel la voiture quitterait le sol.
 - b) Calculer cet angle dans le cas où le conducteur coupe le moteur en A (on prendra $g = 9,8 \text{ m.s}^{-2}$). Conclusion
 - c) Est-il préférable d'accélérer ou de freiner ? Calculer la valeur de $\overline{F}$ pour que la voiture arrive en B sans encombre.

7. Étude d'un skieur, d'après ESIGETEL 2000

On étudie le mouvement d'un skieur descendant une piste selon la ligne de plus grande pente, faisant l'angle α avec l'horizontale. L'air exerce une force de frottement supposée de la forme $\overrightarrow{F} = -\lambda \overrightarrow{v}$, où λ est un coefficient constant positif et $\overrightarrow{v}$ la vitesse du skieur.

On note $\overrightarrow{T}$ et $\overrightarrow{N}$ les composantes tangentielle et normale de la force de frottement exercée par la neige et f le coefficient de frottement solide tel que $\|\overrightarrow{T}\| = f\|\overrightarrow{N}\|$.

On choisit comme origine de l'axe Ox de la ligne de plus grande pente la position initiale du skieur, supposé partir à l'instant initial avec une vitesse négligeable. On note Oy la normale à la piste dirigée vers le haut.

1. Calculer $\overrightarrow{T}$ et $\overrightarrow{N}$.
2. Calculer la vitesse $\overrightarrow{v}$ et la position x du skieur à chaque instant.
3. a) Montrer que le skieur atteint une vitesse limite $\overrightarrow{v_l}$ et calculer $\overrightarrow{v}$ en fonction de $\overrightarrow{v_l}$.
- b) Application numérique : calculer v_l avec $\lambda = 1 \text{ S.I.}$, $f = 0,9$, $g = 10 \text{ m.s}^{-2}$, $m = 80 \text{ kg}$ et $\alpha = 45^\circ$.
4. Calculer littéralement et numériquement la date t_1 où le skieur a une vitesse égale à $v_l/2$.
5. À la date t_1 , le skieur tombe. On néglige alors la résistance de l'air, et on considère que le coefficient de frottement sur le sol est multiplié par 10. Calculer la distance parcourue par le skieur avant de s'arrêter.

8. Atterrissage en catastrophe d'un avion, d'après Agrégation de Chimie 2006

Un avion de chasse de masse $m = 9 \text{ t}$ en panne de freins atterrit à une vitesse $v_a = 241 \text{ km.h}^{-1}$. Une fois au sol, il est freiné en secours par un parachute de diamètre $D = 3 \text{ m}$ déployé instantanément par le pilote au moment où les roues de l'avion touchent le sol. On néglige les forces de frottement des roues sur le sol par rapport à la force de traînée (frottement fluide) du parachute qui s'écrit $T = C_x \rho S v^2 / 2$ avec v la vitesse de l'avion, S la surface projetée du parachute sur un plan perpendiculaire à la vitesse, C_x le coefficient de traînée supposé constant et égal à 1,5, $\rho = 1,2 \text{ kg.m}^{-3}$ la masse volumique de l'air. On considère que le réacteur ne délivre plus aucune poussée.

1. a) Établir l'équation différentielle à laquelle obéit la vitesse v de l'avion.
- b) En déduire l'expression de v en fonction de la date t . On prendra comme origine des temps la date à laquelle les roues de l'avion touchent le sol.
- c) Montrer que cette force de traînée n'est pas suffisante pour stopper complètement l'avion.

2. Dans le cas où les freins fonctionnent, la distance d'atterrissage de l'avion est de l'ordre de $d = 1400$ m. Déterminer la vitesse atteinte par l'avion après avoir été freiné uniquement par le parachute sur cette distance. Faites l'application numérique.

9. Transports de troncs d'arbres, d'après CAPES Agricole 2003

Un exploitant forestier doit évacuer des troncs d'arbres situés sur une parcelle de terrain en montagne en bordure d'un lac. Pour faciliter le transport, il utilise le poids des arbres au cours de leur descente pour en tracter d'autres à travers le lac. Le schéma du dispositif est représenté sur la figure (5.18).

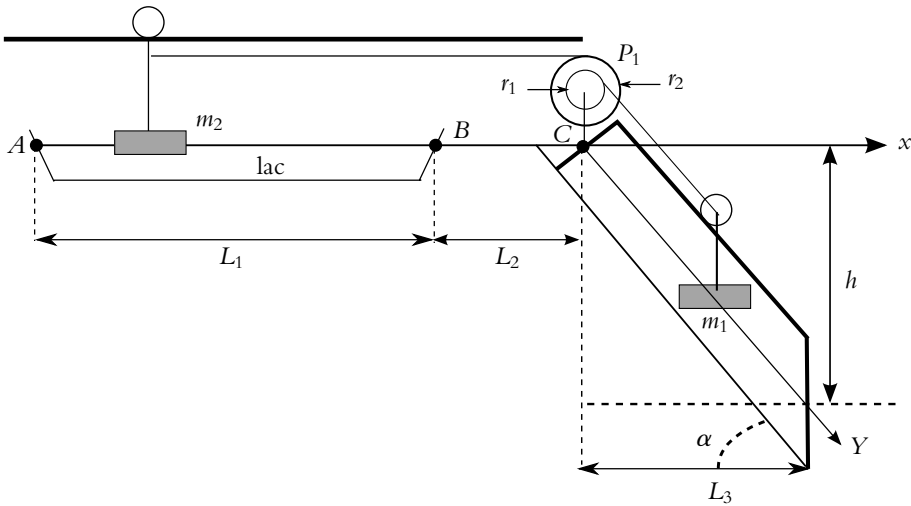


Figure 5.18

Au point C, un portique est installé qui supporte une poulie munie de deux gorges permettant l'enroulement des câbles.

Sur la gorge de rayon r_1 sera enroulé le câble retenant dans sa descente un tronc de masse m_1 assimilé pour l'étude à un point matériel.

Sur la gorge de rayon r_2 sera enroulé la câble tractant deux arbres à travers le lac, symbolisés par une masse ponctuelle m_2 .

Les troncs d'arbres sont suspendus à des câbles guides par des poulies permettant un déplacement considéré sans frottement sur le câble.

Les câbles sont supposés inextensibles et de masse négligeable.

Les troncs se déplaçant dans l'eau, subissent une force de frottement visqueux de la part de l'eau : $\vec{f} = -k\vec{v}$.

Les troncs et leurs support sont assimilés à des points matériels dont la masse totale est celle des troncs.

Données : $g = 9,8 \text{ m.s}^{-2}$; $m_1 = 500 \text{ kg}$; $m_2 = 1000 \text{ kg}$; $h = 300 \text{ m}$; $L_1 = 840 \text{ m}$; $L_2 = 4 \text{ m}$; $L_3 = 300 \text{ m}$; $k = 30 \text{ kg.s}^{-1}$; $r_1 = 20 \text{ cm}$; $r_2 = 40 \text{ cm}$.

1. Déplacement sur l'eau

- a) On repère le déplacement de m_2 par la variable x de l'axe Ax (figure 5.18). Établir l'équation différentielle du mouvement de m_2 . On note T_2 la norme de la tension du câble.
- b) On repère le déplacement de m_1 par la variable Y de l'axe CY parallèle à la ligne de plus grande pente (figure 5.18). On note T_1 la norme de la tension du câble. Établir l'équation différentielle du mouvement de m_1 .
- c) Déterminer le lien entre $\ddot{x}$ et $\ddot{Y}$ dû à la poulie.
- d) On admet que la poulie C impose une relation entre les tensions :

$$T_2 r_2 - T_1 r_1 + \Gamma = 0$$

Γ étant une grandeur positive constante modélisant l'action d'un frein sur la poulie. Dédurre des calculs précédents, l'équation différentielle en $\dot{x}$ en fonction de m_1 , $M = m_2 + m_1 \left(\frac{r_1}{r_2} \right)^2$, Γ , k , g et $\sin \alpha$.

- e) Montrer l'existence d'une vitesse limite v_L et d'une constante de temps τ . Donner leurs expressions littérales et leurs valeurs numériques (on prendra $\Gamma = 680 \text{ N.m}$).
- f) Déterminer l'expression de $\dot{x}(t)$ sachant qu'à $t = 0$, les troncs sont en A et sans vitesse initiale. Évaluer l'ordre de grandeur du temps au bout duquel la vitesse limite est atteinte.
- g) Déterminer l'expression complète $x(t)$.
- h) Calculer le temps t_2 nécessaire aux troncs pour traverser le lac à vitesse constante v_L . Justifier que le temps réel de parcours sur le lac de A à B pourra être assimilé à t_2 .

2. Déplacement sur le sol

Après leur sortie du lac en B , les arbres se déplacent sur le sol. Ils sont alors soumis à un frottement solide. On rappelle que les normes de la composante tangentielle $\vec{R}_t$ et de la composante normale $\vec{R}_n$ de la réaction du sol sont reliées par la relation $R_t = f R_n$, f étant le coefficient de frottement, ici égal à 0,2, R_n et R_t les normes de $\vec{R}_n$ et $\vec{R}_t$. On n'exerce alors plus de freinage sur m_1 ($\Gamma = 0$).

- a) Établir la nouvelle équation différentielle du second ordre en x en fonction de m_1 , m_2 , f , M , g et $\sin \alpha$.
- b) Calculer l'accélération $\ddot{x}$ et en déduire le mouvement de m_2 .
- c) Calculer l'accélération $\ddot{Y}$ de m_1 .
- d) Montrer que la vitesse de m_2 s'annule en un point noté E_2 . Quelle est alors la distance d_2 parcourue sur le sol de B à E_2 ?
- e) Préciser la position de m_1 sur le plan incliné par rapport au point d'arrivée D .

Aspects énergétiques de la dynamique du point

6

L'objet de ce chapitre est de définir les notions énergétiques pour la mécanique du point matériel. On va donc commencer par définir le travail et la puissance d'une force. Les théorèmes de l'énergie cinétique et de l'énergie mécanique pourront alors être établis en introduisant la notion de forces conservatives et d'énergie potentielle.

1. Travail et puissance d'une force

1.1 Position du problème

Soit un point M animé d'une vitesse $\vec{v}$ dans un référentiel $\mathcal{R}$ et soumis à une force $\vec{f}(M)$, dépendant éventuellement des coordonnées du point M (par exemple la tension d'un ressort).

À l'instant t , le point M se trouve en M_t et à l'instant $t+dt$ il est en M_{t+dt} , cet état étant infiniment proche de M_t si dt est suffisamment faible.

On notera $d\vec{OM}$ le déplacement $\overrightarrow{M_t M_{t+dt}}$.

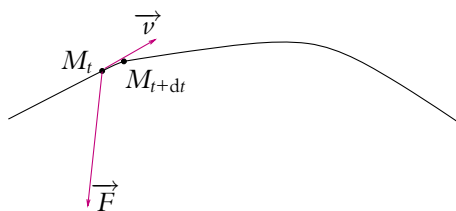


Figure 6.1 Travail d'une force au cours d'un déplacement élémentaire.

1.2 Puissance d'une force

On définit la *puissance* de la force $\vec{f}(M)$ appliquée au point matériel M animé de la vitesse $\vec{v}_{/\mathcal{R}}(M)$ dans un référentiel $\mathcal{R}$ par le produit scalaire :

$$\mathcal{P} = \vec{f}(M) \cdot \vec{v}_{/\mathcal{R}}(M)$$

Comme la vitesse est déterminée par rapport à un référentiel et que la puissance dépend de la vitesse, la puissance est une quantité qui dépend du référentiel dans lequel on travaille.

On distingue trois types de puissance :

- la puissance motrice si $\mathcal{P} > 0$, cela correspond au cas où la force et la projection de la vitesse sur l'axe de la force sont de même sens,
- la puissance résistante si $\mathcal{P} < 0$, cela correspond au cas où la force et la projection de la vitesse sur l'axe de la force sont de sens opposé,
- la puissance nulle si $\mathcal{P} = 0$, cela correspond soit au cas où le point M est immobile (sa vitesse est nulle) soit au cas où le mouvement est perpendiculaire à l'axe de la force (ce sera le cas par exemple de la force magnétique).

1.3 Travail élémentaire d'une force

On appelle *travail élémentaire* d'une force $\vec{f}(M)$ appliquée en un point M au cours du déplacement élémentaire $\overrightarrow{M_t M_{t+dt}}$ pendant l'intervalle de temps dt la quantité :

$$\delta W(M) = \vec{f}(M) \cdot \overrightarrow{M_t M_{t+dt}} = \vec{f}(M) \cdot d\overrightarrow{OM}$$

On note δW et non pas dW cette quantité car le travail ne correspond pas à la différence d'une grandeur entre deux états mais à une variation au cours d'un déplacement. Le travail sera très fortement dépendant du déplacement : le plus souvent, la valeur du travail au cours du déplacement entre deux points dépend du chemin suivi. Cette notion sera revue dans le cours de thermodynamique.

On peut également définir le travail élémentaire à partir de la puissance $\mathcal{P}$. En effet, par définition de la vitesse, on a :

$$d\overrightarrow{OM} = \vec{v}_{/\mathcal{R}}(M)dt$$

ce qui permet d'écrire :

$$\delta W(M) = \vec{f}(M) \cdot \vec{v}_{/\mathcal{R}}(M)dt = \mathcal{P}dt$$

La puissance est donc le travail de la force par unité de temps :

$$\mathcal{P} = \frac{\delta W}{dt}$$

1.4 Travail d'une force au cours d'un déplacement

Au cours d'un déplacement le long d'un chemin Γ_{AB} , le travail de la force correspond à la somme des travaux élémentaires, ce qui se traduit par l'intégrale du travail élémentaire :

$$W_{A \rightarrow B} = \int_{M \in \Gamma_{AB}} \delta W(M) = \int_{M \in \Gamma_{AB}} \vec{f}(M) \cdot d\overrightarrow{OM} = \int_{t_A}^{t_B} \mathcal{P}(t)dt$$

à savoir également l'intégrale de la puissance de la force au cours du déplacement, puissance que l'on peut considérer comme une fonction du temps par l'intermédiaire des coordonnées du point M et du vecteur-vitesse.



Le travail comme la puissance dépend du référentiel dans lequel on se place.

2. Théorème de l'énergie cinétique

2.1 Définition de l'énergie cinétique

On appelle *énergie cinétique* la quantité scalaire :

$$Ec_{/\mathcal{R}}(M) = \frac{1}{2}mv_{/\mathcal{R}}^2(M)$$

pour une particule ponctuelle de masse m animée d'une vitesse $\vec{v}_{/\mathcal{R}}(M)$ dans le référentiel $\mathcal{R}$ considéré. Par la suite, on ne précisera plus par rapport à quel référentiel l'énergie cinétique est définie mais on n'oubliera pas pour autant qu'il s'agit, comme la puissance ou le travail, d'une quantité qui dépend du référentiel dans lequel on la calcule, par l'intermédiaire de la vitesse du point.

2.2 Théorème de l'énergie cinétique en référentiel galiléen

Ce théorème s'énonce de la façon suivante :

la variation d'énergie cinétique d'un point matériel entre deux instants est égale au travail des forces qui s'exercent sur ce point entre les deux instants considérés.

Il s'agit d'une conséquence du principe fondamental de la dynamique. En effet, celui-ci s'écrit (pour simplifier la vitesse du point M dans le référentiel $\mathcal{R}$ est notée $\vec{v}$) :

$$m \frac{d\vec{v}}{dt} = \sum_i \vec{f}_i$$

En multipliant scalairement cette relation par $\vec{v} dt$, on obtient :

$$m \frac{d\vec{v}}{dt} \cdot \vec{v} dt = \left(\sum_i \vec{f}_i \right) \cdot \vec{v} dt = \sum_i \left(\vec{f}_i \cdot \vec{v} \right) dt = \sum_i \delta W_i = \delta W$$

en notant δW_i le travail de la force $\vec{f}_i$ et δW la somme des différents travaux. Or

$$m \frac{d\vec{v}}{dt} \cdot \vec{v} dt = \frac{d}{dt} \left(\frac{1}{2}mv^2 \right) dt = d \left(\frac{1}{2}mv^2 \right)$$

soit finalement :

$$dEc = d \left(\frac{1}{2}mv^2 \right) = \delta W$$

On en déduit par intégration le travail de la somme des forces entre deux instants t_1 et t_2 :

$$W_{t_1 \rightarrow t_2} = \frac{1}{2}mv^2(t_2) - \frac{1}{2}mv^2(t_1) = Ec(t_2) - Ec(t_1) = \Delta Ec$$

Le théorème de l'énergie cinétique peut se formuler de trois manières :

1. expression intégrale :

$$Ec(t_2) - Ec(t_1) = \frac{1}{2}mv^2(t_2) - \frac{1}{2}mv^2(t_1) = W_{t_1 \rightarrow t_2}(\vec{f})$$

2. expression différentielle :

$$dEc = \delta W(\vec{f})$$

3. expression en fonction de la puissance :

$$\frac{dEc}{dt} = \mathcal{P}(\vec{f})$$

2.3 Utilisation du théorème de l'énergie cinétique

L'utilisation du théorème de l'énergie cinétique est de deux sortes :

- on déduit le travail de la somme des forces appliquées à partir de la connaissance de la variation d'énergie cinétique,
- on déduit la variation de l'énergie cinétique connaissant le travail de la somme des forces.

Cette deuxième utilisation est souvent délicate car le travail des forces ne s'exprime pas de façon simple en général. En revanche, il est particulièrement efficace dans le cas où l'expression simple du travail est connue comme le montre l'exemple du paragraphe 2.5.

2.4 Intérêt d'une approche énergétique pour les systèmes à un degré de liberté

Le théorème de l'énergie cinétique qui vient d'être établi est une conséquence du principe fondamental de la dynamique et n'introduit pas de nouveau postulat. Une approche énergétique revient donc implicitement à utiliser le principe fondamental de la dynamique.

En revanche, l'expression du principe fondamental de la dynamique est vectorielle et correspond à trois équations scalaires dans l'espace tandis que le théorème de l'énergie cinétique ne fournit qu'une équation scalaire. En passant du principe fondamental de la dynamique au théorème de l'énergie cinétique, on a donc perdu deux équations

scalaires. Lorsqu'on traite un problème à un degré de liberté c'est-à-dire dont la résolution ne concerne qu'une variable, cela ne pose aucune difficulté puisqu'une seule équation suffit pour déterminer une seule variable. Par contre, si la situation nécessite de connaître l'expression de deux variables ou plus, la seule utilisation du théorème de l'énergie cinétique ne permettra pas la résolution complète du problème. On devra alors recourir à des projections du principe fondamental de la dynamique sur des axes différents pour compléter la formalisation énergétique.



On retiendra donc qu'une méthode énergétique est adaptée aux cas où un seul paramètre de position suffit à décrire l'évolution du système à savoir aux systèmes à un degré de liberté.

2.5 Exemple d'étude d'un problème à l'aide du théorème de l'énergie cinétique

On considère une bille assimilée à un point matériel de masse m pouvant se déplacer sans frottement sur une demi-sphère de centre O et de rayon a . On l'abandonne sans vitesse initiale depuis le sommet de la sphère et on cherche à déterminer les conditions nécessaires pour que la bille ne décolle pas.

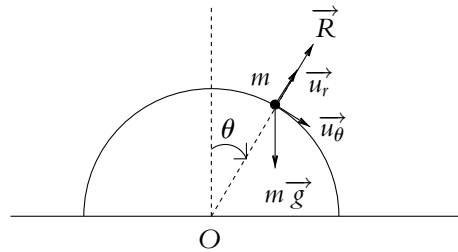


Figure 6.2 Point matériel sur une demi-sphère.

1. Système : masse m .
2. Référentiel du laboratoire considéré comme galiléen.
3. Bilan des forces :
 - poids $m\vec{g}$
 - réaction de la sphère perpendiculaire à la surface du fait de l'absence de frottement.
4. Résolution par le théorème de l'énergie cinétique.

On utilise les coordonnées polaires dans le plan vertical où a lieu le mouvement ; la vitesse s'écrit : $\vec{v} = a\dot{\theta}\vec{u}_\theta$. L'énergie cinétique est :

$$E_c = \frac{1}{2}ma^2\dot{\theta}^2$$

La réaction, étant perpendiculaire au mouvement, ne travaille pas et le travail du poids entre l'instant $t = 0$ et un instant t où la position du point $\square$ est repérée par l'angle θ vaut :

$$W_{\text{poids}} = \int_0^t m\vec{g} \cdot \vec{v} dt = \int_0^t mga\dot{\theta} \sin \theta dt = \int_0^\theta mga \sin \theta d\theta = mga(1 - \cos \theta)$$

Le théorème de l'énergie cinétique entre ces deux instants donne :

$$\frac{1}{2}ma^2\dot{\theta}^2 - 0 = mga(1 - \cos \theta)$$

soit

$$\dot{\theta}^2 = \frac{2g}{a}(1 - \cos \theta)$$

Le principe fondamental de la dynamique s'écrit : $m\vec{a} = m\vec{g} + \vec{R}$. On souhaite l'expression de la réaction donc on projette cette relation sur $\vec{u}_r$. On obtient : $R - mg \cos \theta = -ma\dot{\theta}^2$ soit, en remplaçant $\dot{\theta}^2$ par l'expression établie précédemment

$$R = mg(3 \cos \theta - 2)$$

La bille reste en contact avec la sphère tant que la réaction ne s'annule pas. Il y aura décollage pour $R = 0$ donc pour θ_0 tel que

$$\cos \theta_0 = \frac{2}{3} \quad \text{ou} \quad \text{encore} \quad \theta_0 = 48,2^\circ$$

3. Énergie potentielle et forces conservatives

3.1 Définitions

On dit qu'une force *dérive d'un potentiel* ou encore qu'elle est *conservative* si le travail entre deux points A et B ne dépend pas du chemin suivi mais uniquement des points A et B . On peut alors écrire le travail W_{AB} entre A et B comme la différence $Ep(A) - Ep(B)$ où Ep est une fonction de la variable de position. Au niveau élémentaire, la relation s'écrit : $\delta W = -dEp$. On appellera *énergie potentielle de la force* cette fonction Ep .

L'énergie potentielle étant définie à partir de sa variation, elle n'est déterminée qu'à une constante près.

3.2 Interprétation physique

La dénomination de force conservative se justifie à partir du théorème de l'énergie cinétique appliqué à ce type de forces.

Si un point matériel n'est soumis qu'à une force conservative, le travail de la force entre les instants t_1 où le point est en A et t_2 où il est en B s'écrit : $W = Ep(A) - Ep(B) = -\Delta Ep$. Alors le théorème de l'énergie cinétique s'écrit :

$$\Delta Ec = W = -\Delta Ep$$

où $\Delta Ec = Ec(t_2) - Ec(t_1)$

soit

$$\Delta(E_c + E_p) = 0$$

Cette équation traduit la conservation de la quantité $E_c + E_p$ à savoir la somme de l'énergie cinétique et de l'énergie potentielle. Ceci explique le nom donné à ce type de force.

3.3 Exemples de forces conservatives

a) Poids d'un corps

Le poids d'un point matériel de masse m dans le champ de pesanteur $\vec{g}$ vaut : $m\vec{g}$, soit en projection dans un système de coordonnées cartésiennes dont l'axe Oz est la verticale dirigée vers le haut : $-mg\vec{u}_z$. Le travail de cette force s'écrit donc :

$$\delta W = m\vec{g} \cdot d\vec{OM} = -mgdz = -d(mgz)$$

donc

$$E_p = mgz + C$$

où C est une constante.



On n'oubliera pas que dans cette expression, l'axe Oz est dirigé selon la verticale ascendante.

b) Force électrostatique

Soit un point matériel de charge électrique q placé dans un champ électrique uniforme $\vec{E}$. Il est soumis à la force

$$\vec{F} = q\vec{E}$$

dont le travail élémentaire s'écrit :

$$\delta W = q\vec{E} \cdot d\vec{OM}$$

On établira dans le cours d'électrostatique l'existence du potentiel électrostatique V tel que :

$$\vec{E} \cdot d\vec{OM} = -dV$$

Cela permet d'écrire le travail de la force électrostatique sous la forme :

$$\delta W = -q dV$$

et d'en déduire l'énergie potentielle correspondante

$$E_p = qV + C$$

où C est une constante.

c) Force de rappel d'un ressort

Un ressort de raideur k ayant subi un allongement $\Delta l = l - l_0$ exerce une force de rappel dans la direction de l'allongement (donnée par le vecteur unitaire $\vec{u}_{ext}$) :

$$\vec{f} = -k\Delta l \vec{u}_{ext}$$

Pour un déplacement le long de cette direction, on a : $d\vec{OM} = dl \vec{u}_{ext}$. Le calcul du travail de cette force donne :

$$\delta W = -k\Delta l \vec{u}_{ext} \cdot dl \vec{u}_{ext} = -k(l - l_0) dl = -d\left(\frac{1}{2}k(l - l_0)^2\right)$$

ce qui conduit à l'expression de l'énergie potentielle :

$$Ep = \frac{1}{2}k(l - l_0)^2 + C$$

où C est une constante.

Si le déplacement a lieu selon une autre direction, on peut l'écrire :

$$d\vec{OM} = dl \vec{u}_{ext} + ld \vec{u}_{ext}$$

Or $\vec{u}_{ext}$ est un vecteur unitaire soit $\vec{u}_{ext} \cdot \vec{u}_{ext} = 1$ donc $\vec{u}_{ext} \cdot d\vec{u}_{ext} = 0$. On en déduit le travail élémentaire :

$$\delta W = -k(l - l_0) dl$$

On aura donc toujours la même expression de l'énergie potentielle élastique pour la force de rappel du ressort :

$$Ep = \frac{1}{2}k(l - l_0)^2 + C$$

3.4 Exemples de forces non conservatives

Si on considère une force de frottement fluide par exemple proportionnelle à la vitesse, on a : $\vec{f} = -\lambda \vec{v}$ et le travail élémentaire de cette force s'écrit :

$$\delta W = \vec{f} \cdot d\vec{OM} = -\lambda \vec{v} \cdot \vec{v} dt = -\lambda v^2 dt$$

On ne peut pas l'écrire sous la forme d'une différentielle : il ne s'agit donc pas d'une force conservative.



On remarque que le travail est toujours négatif. Si la force était conservative, on pourrait écrire :

$$W_{AB} = Ep(A) - Ep(B) = -W_{BA} \quad (6.1)$$

Comme ici le travail est toujours négatif, $W_{AB} < 0$ et $W_{BA} < 0$, la relation (6.1) ne peut donc pas être vérifiée. La force n'est pas conservative.

4. Énergie mécanique

4.1 Définition de l'énergie mécanique

Dans le cas où le système n'est soumis qu'à des forces conservatives ou à des forces qui ne travaillent pas, on a établi que :

$$W = -\Delta E_p = \Delta E_c$$

soit

$$\Delta (E_c + E_p) = 0$$

On en déduit :

$$E_c + E_p = E_m$$

où E_m est une constante.

La quantité E_m qui est la somme de l'énergie cinétique et des énergies potentielles est appelée *énergie mécanique* du point matériel.

4.2 Conservation de l'énergie mécanique

D'après les résultats du paragraphe précédent, l'énergie mécanique d'un point matériel soumis uniquement à des forces conservatives ou à des forces qui ne travaillent pas est une constante du mouvement.

On parle d'*intégrale première du mouvement*. En effet, est appelée intégrale première du mouvement toute quantité qui se conserve au cours du mouvement et qui n'est fonction que de la position et de ses dérivées premières par rapport au temps.

On obtient donc une loi de conservation quand les forces dérivent d'un potentiel. L'énergie potentielle peut varier en se transformant en énergie cinétique et réciproquement l'énergie cinétique peut se transformer en énergie potentielle, la somme des deux restant constante.

Ce résultat n'est bien sûr valable que lorsque toutes les forces sont conservatives ou qu'elles ne travaillent pas.

4.3 Cas général : non conservation de l'énergie mécanique

Dans le cas général, il est nécessaire de distinguer les forces conservatives de celles qui ne le sont pas. En notant leurs travaux respectifs W et W' , on sait d'après ce qui précède que W peut se mettre sous la forme : $W = -\Delta E_p$. En revanche, il n'est pas aisé d'avoir une expression simple de W' . On en déduit donc l'expression du travail élémentaire total :

$$W_{tot} = W + W' = -\Delta E_p + W'$$

En appliquant le théorème de l'énergie cinétique à cette situation, on a donc :

$$\Delta E_c = W_{tot} = -\Delta E_p + W'$$

ce qui permet d'en déduire :

$$\Delta E_m = \Delta (E_c + E_p) = W'$$

ou sous forme élémentaire :

$$dE_m = d(E_c + E_p) = \delta W'$$

Parfois ce résultat est appelé « théorème de l'énergie mécanique » que l'on peut formuler de la manière suivante :

la variation d'énergie mécanique au cours du mouvement est égale au travail des forces qui ne dérivent pas d'un potentiel, autrement dit des forces non conservatives.

Ce n'est qu'une autre formulation du théorème de l'énergie cinétique.

Il est à noter que le problème de la conservation de l'énergie fait partie d'un cadre beaucoup plus général que celui de la mécanique. On reviendra sur ce point dans le cours de thermodynamique et plus précisément lors de l'étude du premier principe qui postule la conservation de l'énergie totale. Le résultat qui vient d'être établi sera alors généralisé.

5. Application à l'étude qualitative des mouvements et des équilibres

Dans ce paragraphe, on se place dans le cas où l'énergie mécanique se conserve c'est-à-dire dans le cas où toutes les forces qui travaillent sont conservatives.

5.1 Conséquence de la positivité de l'énergie cinétique

Ici, la conservation de l'énergie mécanique s'écrit :

$$E_c + E_p = E_m = \text{constante}$$

On en déduit :

$$E_c = \frac{1}{2}mv^2 = E_m - E_p$$

Par nature, E_c est donc une quantité positive donc **l'énergie mécanique E_m doit être supérieure à l'énergie potentielle E_p** pour qu'il y ait mouvement.

5.2 Nature du mouvement en fonction de l'énergie mécanique

La connaissance de l'énergie potentielle E_p en fonction de la variable de position qu'on notera x dans la suite permet de distinguer les différents mouvements possibles.

Considérons par exemple l'énergie potentielle décrite par la figure ci-contre.

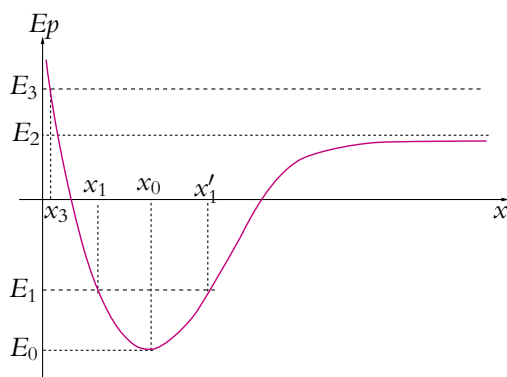


Figure 6.3 Exemple d'énergie potentielle.

- Si $E_m < E_0$, alors $E_m < E_p$ quelle que soit la position et il ne peut y avoir de mouvement d'après le paragraphe précédent.
- Si $E_m = E_0$, le point matériel est en $x = x_0$ et on a un état stationnaire (x reste égal à x_0).
- Si $E_0 < E_m < E_2$ (si $E_m = E_1$ dans l'exemple considéré) alors le point matériel ne peut évoluer qu'entre les positions x_1 et x'_1 . On dit que le point matériel est dans un *état lié*, à savoir que le point reste dans une zone finie de l'espace et ne part pas à l'infini.
- Si $E_m \geq E_2$ (si $E_m = E_3$ dans l'exemple considéré) alors les seuls mouvements possibles sont ceux pour lesquelles le paramètre de position vérifie $x \geq x_3$ afin d'assurer la condition $E_m \geq E_p$. Le point matériel peut atteindre la position $x \rightarrow \infty$, on dit qu'il est dans un *état libre*.

5.3 Équilibre

Un point matériel est en équilibre lorsque sa vitesse et son accélération sont nulles.

D'après le principe fondamental de la dynamique, la somme des forces qui s'appliquent sur ce point est nulle :

$$\sum_i \vec{f}_i = \vec{0}$$

5.4 Détermination des positions d'équilibre

On cherche à déterminer les positions pour lesquelles le point matériel est à l'équilibre.

Lorsque l'ensemble des forces qui s'exercent sur un point matériel en équilibre sont conservatives, on peut utiliser un raisonnement énergétique. On désigne par E_p l'énergie potentielle totale et on suppose qu'elle ne dépend que d'une variable de position notée x .

Le travail des forces s'écrit :

$$\delta W = \vec{F} \cdot d\vec{OM} = F_x dx = -dEp$$

donc :

$$F_x = -\frac{dEp}{dx}$$

La condition d'équilibre impose alors d'avoir :

$$\frac{dEp}{dx} = 0$$

L'énergie potentielle passe par un extremum (minimum ou maximum) au(x) point(s) d'équilibre.

5.5 Étude de la stabilité des équilibres

a) Définition

On dit qu'un équilibre est *stable* si quand on l'écarte de sa position d'équilibre, le point matériel revient vers sa position d'équilibre initiale.

Un équilibre sera dit *instable* si le fait de l'écarter de sa position d'équilibre éloigne définitivement le point de cette position.

b) Analyse de la stabilité pour les forces conservatives

Au voisinage d'une position d'équilibre x_0 , il est possible d'effectuer un développement limité de l'énergie potentielle à l'ordre 2 :

$$Ep(x) = Ep(x_0) + \left(\frac{dEp}{dx} \right)_{x=x_0} (x - x_0) + \left(\frac{d^2Ep}{dx^2} \right)_{x=x_0} \frac{(x - x_0)^2}{2}$$

Or x_0 est une position d'équilibre donc $\left(\frac{dEp}{dx} \right)_{x=x_0} = 0$. Alors

$$Ep(x) = Ep(x_0) + \left(\frac{d^2Ep}{dx^2} \right)_{x=x_0} \frac{(x - x_0)^2}{2}$$

On en déduit que le fait d'écarter légèrement de dx le point de sa position d'équilibre consiste à exercer une force élémentaire dF_x telle que

$$dF_x = - \left(\frac{dEp}{dx} \right) (x = x_0 + dx) = - \left(\frac{d^2Ep}{dx^2} \right)_{x=x_0} dx$$

Si on exerce un déplacement $dx > 0$, il faut que $dF_x < 0$ pour ramener le point à sa position initiale. Cela impose donc :

$$\frac{d^2Ep}{dx^2}(x = x_0) > 0$$

On retrouve le même résultat pour un déplacement $dx < 0$ (il faut alors que $dF_x(x = x_0) > 0$). On en déduit qu'on aura un **équilibre stable en x_0 si l'énergie potentielle est minimale en x_0** .

A contrario, l'équilibre en x_0 sera instable si l'énergie potentielle est maximale en x_0 .

Si $\frac{d^2Ep}{dx^2} = 0$, on dit que l'équilibre est indifférent : le point peut ou non revenir à sa position d'équilibre. Il faudrait étudier les dérivées d'ordre supérieur pour conclure.

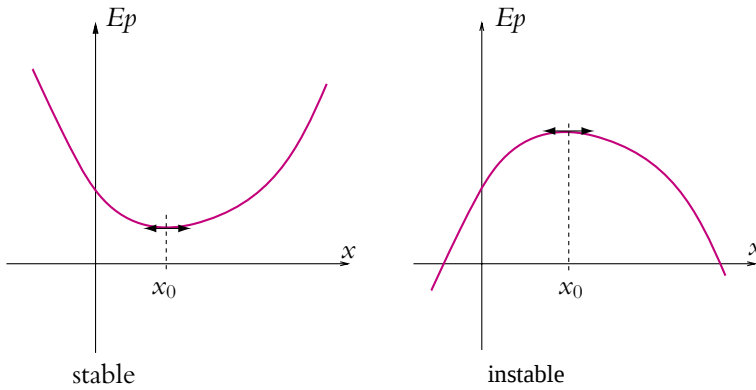


Figure 6.4 Énergie potentielle et stabilité des positions d'équilibre.

6. Approche à l'aide des portraits de phase

L'objet de ce paragraphe est de présenter une représentation graphique appelée *portrait de phase* qui permet l'analyse qualitative d'un système. On se limite aux systèmes à un degré de liberté c'est-à-dire aux systèmes dont l'évolution est décrite par un seul paramètre de position. Ce paramètre sera noté x dans la suite de l'exposé.

6.1 Définition

Dans les exemples étudiés précédemment, on a obtenu différents types d'équations du mouvement :

- $\ddot{x} = -g$ pour la chute libre d'un point matériel en l'absence de résistance de l'air,
- $\ddot{x} + \frac{k}{m}\dot{x} = -g$ pour la chute libre d'un point matériel en tenant compte de la résistance de l'air qu'on suppose proportionnelle à la vitesse,
- $\ddot{x} + \frac{k}{m}x = 0$ pour le mouvement d'une masse suspendue à un ressort dont l'autre extrémité est fixe,
- etc.

Dans tous les cas, il est possible de mettre ces équations sous la forme :

$$\ddot{x} = f(\dot{x}, x) \quad (6.2)$$

qui définit un système dynamique à deux degrés de liberté c'est-à-dire dont l'évolution dépend de deux paramètres : la position x et la vitesse $\dot{x}$ du système.

On peut noter qu'ici on travaille avec un problème unidimensionnel, ce qui permet de n'avoir qu'une composante de la position et une composante de la vitesse. Il est important de remarquer que la dimension du problème (ici un) est différente de la notion de degrés de liberté qui donne le nombre de variables nécessaires pour décrire la totalité du système (ici l'équation différentielle est du second ordre, il y a donc deux degrés de liberté). Du fait de la dimension du problème, une seule composante d'espace est nécessaire. Comme on a deux degrés de liberté, il faudra deux variables pour décrire l'état du système soit par exemple la position et une de ses dérivées par rapport au temps.

On peut remplacer l'équation différentielle (6.2) du second ordre à une variable par le système de deux équations différentielles du premier ordre à deux variables :

$$\begin{cases} v = \frac{dx}{dt} \\ \frac{dv}{dt} = f(x, v) \end{cases}$$

ce qui permet d'avoir accès à la position $x(t)$ et à la vitesse $v(t)$.

6.2 Solution et conditions initiales

La donnée de conditions initiales (une position initiale x_0 et une vitesse initiale v_0) permet en général d'obtenir une solution unique aux équations précédentes. Il suffit pour cela que la fonction f soit suffisamment régulière pour permettre les intégrations nécessaires.

Une unique solution est obtenue pour chaque condition initiale qui se traduit par une trajectoire. On a la connaissance complète du système : position, vitesse et accélération. Il s'agit de la traduction du déterminisme propre à la mécanique classique.

Exemple de la chute libre sans résistance de l'air avec vitesse initiale

La solution générale est :

$$x(t) = -\frac{1}{2}gt^2 + At + B$$

Les conditions initiales $x(0) = x_0$ et $v(0) = \dot{x}(0) = v_0$ permettent d'obtenir $A = v_0$ et $B = x_0$. On a alors une solution unique :

$$x(t) = -\frac{1}{2}gt^2 + v_0t + x_0$$

et

$$v(t) = -gt + v_0$$

6.3 Portrait de phase

On appelle *plan de phase* le plan (x, v) où on porte les valeurs des positions en abscisse et celles de la vitesse en ordonnées.

Dans le plan de phase, les conditions initiales, à savoir ici position initiale x_0 et vitesse initiale v_0 , correspondent à un point M_0 de composantes (x_0, v_0) . La solution unique de l'équation du mouvement correspondant à ces conditions initiales décrit une courbe paramétrée par le temps $(x(t), v(t))$: le point donnant l'état du système à un instant t évolue au cours du temps. La courbe ainsi décrite est appelée *trajectoire de phase*. Elle passe par le point représentant les conditions initiales M_0 .

Pour chaque couple de conditions initiales, on peut effectuer la même étude : on obtient différentes courbes associées chacune à une condition initiale. L'ensemble des courbes associées à l'équation du mouvement d'un système définit le *portrait de phase* du système. Ce dernier sera donc obtenu en traçant les trajectoires de phase pour plusieurs conditions initiales.

6.4 Propriétés des plans de phase

a) Sens de parcours des trajectoires de phase

Les trajectoires de phase sont parcourues dans le sens des aiguilles d'une montre.

En effet, dans le demi-plan supérieur où l'ordonnée est positive, la vitesse est positive. Or la vitesse étant la dérivée de la position par rapport au temps, la fonction donnant la position en fonction du temps $x(t)$ est donc croissante. Si x augmente alors on se déplace de gauche à droite le long de la courbe.

Inversement, dans le demi-plan inférieur où l'ordonnée est négative, la vitesse est négative et la fonction donnant la position en fonction du temps $x(t)$ est donc décroissante. Si x diminue alors on se déplace de droite à gauche le long de la courbe.

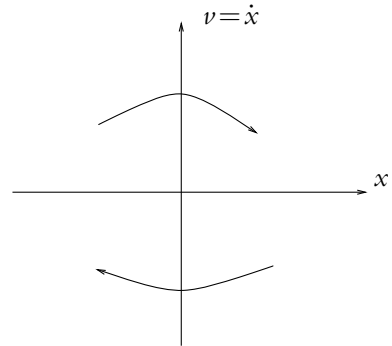


Figure 6.5 Sens de parcours des trajectoires de phase.

b) Intersection de l'axe Ox avec les trajectoires de phase

Lorsque la trajectoire de phase coupe l'axe des abscisses Ox , soit elle passe par un point singulier (que l'on va définir) soit elle est perpendiculaire à l'axe Ox .

En effet, la tangente à la trajectoire de phase en un point (x, v) peut être définie comme la droite passant par ce point et faisant un angle α avec l'axe Ox tel que

$$\tan \alpha = \frac{dv}{dx} = \frac{\frac{dv}{dt}}{\frac{dx}{dt}}$$

Or $\frac{d\nu}{dt} = f(x, \nu)$ par définition de l'équation du mouvement et $\frac{dx}{dt} = \nu$ par définition de la vitesse. On en déduit que

$$\tan \alpha = \frac{f(x, \nu)}{\nu}$$

Or lorsque les trajectoires de phase coupent l'axe des abscisses Ox du plan de phase, la vitesse ν est nulle. On distingue alors deux cas :

- si $f(x, \nu) \neq 0$ alors $\tan \alpha \rightarrow \infty$ et la tangente à la trajectoire de phase est une droite parallèle à l'axe des ordonnées : la trajectoire de phase est perpendiculaire à l'axe des abscisses,
- si $f(x, \nu) = 0$ alors $\tan \alpha = 0$ et on obtient une forme indéterminée pour $\tan \alpha$: on ne peut conclure sur la valeur de cette quantité et on dit qu'on a un *point singulier*.

c) Positions d'équilibre dans le plan de phase

Les positions d'équilibre sont définies par :

$$\nu = \dot{x} = 0 \quad \text{et} \quad \ddot{x} = \dot{\nu} = 0$$

soit :

$$\begin{cases} \nu = 0 \\ f(x, \nu) = 0 \end{cases} \quad \text{ou} \quad f(x, 0) = 0$$

Il s'agit des points de l'axe des abscisses où la vitesse est nulle ou encore des points pour lesquels on obtenait une forme indéterminée au paragraphe précédent : les points singuliers.

d) Non-croisement de deux trajectoires de phase

Deux trajectoires de phase correspondent à deux conditions initiales distinctes et ne peuvent se couper en dehors de l'axe des abscisses que si elles sont confondues.

En effet, si on suppose que deux trajectoires de phase se coupent en un point M_0 qu'on peut considérer comme conditions initiales. L'unicité de la solution pour des conditions initiales données prouve que les deux trajectoires sont confondues.

On remarquera cependant que deux trajectoires de phase peuvent se croiser en un point singulier.

e) Trajectoire de phase et systèmes périodiques

Si la trajectoire se referme sur elle-même, on retrouve les conditions initiales et elle se reproduit par la suite identique à elle-même. Le mouvement est périodique et le temps de parcours de la trajectoire de phase est égal à la période du mouvement.

6.5 Exemple de portrait de phase : la chute libre

On a vu que l'équation du mouvement de la chute libre est $\ddot{x} = -g$ et que, par intégration, on a : $v = \dot{x} = -gt + v_0$ et $x = -\frac{1}{2}gt^2 + v_0t + x_0$. En éliminant t de ces deux dernières relations, on obtient l'équation de la trajectoire de phase correspondant aux conditions initiales $x(0) = x_0$ et $v(0) = v_0$:

$$x = -\frac{v^2}{2g} + x_0 + \frac{v_0^2}{2g}$$

Les trajectoires de phase sont donc des paraboles d'axe Ox .

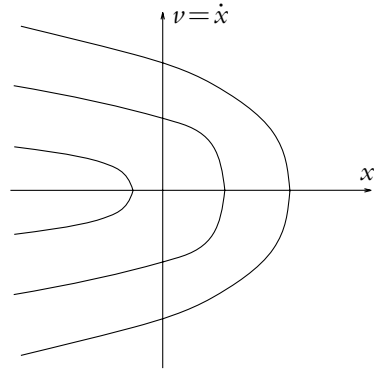


Figure 6.6 Portrait de phase de la chute libre.

6.6 Portraits de phase des systèmes conservatifs

Comme il n'y a que des forces conservatives, l'énergie mécanique se conserve et on peut écrire :

$$Ec + Ep(x) = \frac{1}{2}mv^2 + Ep(x) = Em = \text{constante}$$

Il s'agit bien de l'équation d'une trajectoire de phase puisqu'elle relie la position x et la vitesse $v = \dot{x}$. On en déduit que les trajectoires de phase sont des courbes isoénergétiques (c'est-à-dire ayant même énergie mécanique) pour un système conservatif.

Soit un système dont l'énergie potentielle a l'allure suivante :

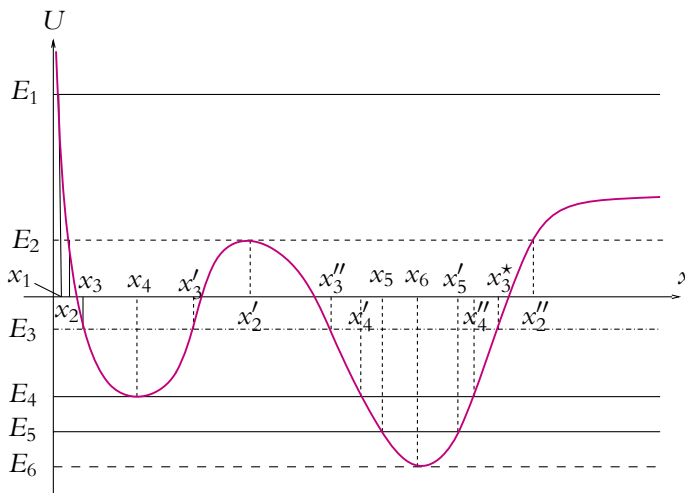


Figure 6.7 Exemple d'énergie potentielle.

Suivant les conditions initiales et la valeur correspondante de l'énergie mécanique, les trajectoires de phase ont différentes allures. On présente dans ce qui suit l'analyse de quelques cas correspondant tous à une vitesse initiale nulle et différentes positions initiales x_0 . On note E_m l'énergie mécanique.

- Pour une énergie mécanique $E_m = E_6$, on a une seule position possible correspondant à une position d'équilibre stable, le point matériel ne bouge pas, la trajectoire de phase se réduit au point d'abscisse $x = x_6$ de l'axe des abscisses.
- Pour une énergie mécanique $E_m = E_5$ et une position initiale $x_0 = x_5$, on a un mouvement oscillatoire entre les points d'abscisse x_5 et x'_5 autour de la position d'équilibre $x = x_6$, la trajectoire de phase est fermée et parcourue dans le sens des aiguilles d'une montre. Le mouvement est périodique.
- Pour une énergie mécanique $E_m = E_4$, on distingue deux cas suivant la position initiale :
 - $x_0 = x_4$: on est dans une situation analogue à celle du cas $x_0 = x_6$ et $E_m = E_6$, le point ne quitte pas sa position d'équilibre stable, la trajectoire de phase est le point d'abscisse $x = x_4$ de l'axe des abscisses.
 - $x_0 = x'_4$: on a un mouvement oscillatoire entre les points d'abscisse x'_4 et x''_4 autour de la position d'équilibre $x = x_6$, la trajectoire de phase est fermée et parcourue dans le sens des aiguilles d'une montre, le mouvement est périodique.
- Pour une énergie mécanique $E_m = E_3$, on a deux mouvements oscillatoires possibles sans qu'un passage de l'un à l'autre puisse avoir lieu : si $x_0 = x_3$, le point oscille entre les points d'abscisse x_3 et x'_3 autour de la position d'équilibre $x = x_4$ et si $x_0 = x''_3$, il oscille entre les points d'abscisse x''_3 et x^*_3 autour de la position d'équilibre $x = x_6$; dans les deux cas, la trajectoire de phase est fermée et parcourue dans le sens des aiguilles d'une montre, le mouvement est périodique.
- Pour une énergie mécanique $E_m = E_2$, on atteint le cas limite de la situation précédente $E_m = E_3$: les deux trajectoires de phase ont un point commun en $x = x'_2$. Il ne s'agit pas d'un point où le système a le choix entre les deux trajectoires : lorsqu'il est sur une trajectoire (soit entre x_2 et x'_2 autour de x_4 soit entre x'_2 et x''_2 autour de x_6), il y reste. Le point $x = x'_2$ est un point d'arrêt. On appelle *séparatrices* les deux trajectoires correspondant à cette situation : elles sont la limite des cas du type $E_m = E_3$ et des cas qui sont décrits après pour $E_m = E_1$. On notera que le point d'arrêt correspond à la position d'équilibre instable.
- Pour une énergie mécanique $E_m = E_1$, on a une trajectoire libre entre x_1 et $+\infty$ autour des trois positions d'équilibre x_4 , x'_2 et x_6 .

Le portrait de phase a l'allure suivante :

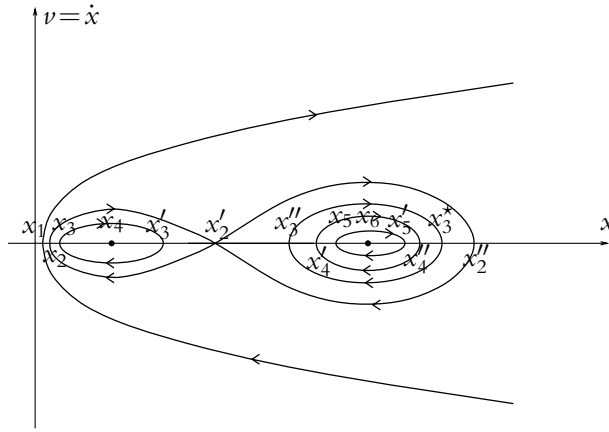


Figure 6.8 Portrait de phase associé à l'énergie potentielle précédente.

Si E_5 est proche de E_6 , les mouvements sont petits autour de la position d'équilibre : on verra au chapitre suivant qu'on retrouve un oscillateur harmonique et que la trajectoire de phase est une ellipse.

Si E_3 est assez grande par rapport à E_4 (donc à E_6), les mouvements sont périodiques mais pas harmoniques et les trajectoires de phase ne sont plus des ellipses.

6.7 Influence des frottements

L'application du théorème de l'énergie cinétique au problème d'un point matériel soumis à une force conservative d'énergie potentielle E_p et à une force de frottement donne :

$$dEm = d(Ec + E_p) = \delta W' < 0$$

donc l'énergie mécanique diminue.

Par conséquent, dans le cas d'un puits de potentiel représenté sur la figure ci-contre, la diminution de l'énergie mécanique conduit à une réduction de l'amplitude des oscillations du système et au bout d'un temps infini, le système se retrouve

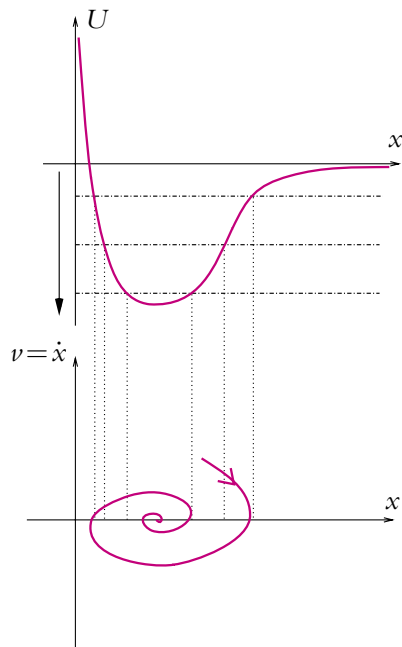


Figure 6.9 Influence des frottements sur un portrait de phase.

dans sa position d'équilibre stable. On en déduit l'allure de la trajectoire de phase qui se rapproche petit à petit de la position d'équilibre.

Cas d'un système amplificateur

Dans le cas d'un système amplificateur, on fournit de l'énergie au système : on obtient donc l'effet inverse, la trajectoire de phase s'éloigne de plus en plus de la position d'équilibre.

7. Étude énergétique d'une masse suspendue à un ressort

On reprend l'exemple d'une masse ponctuelle suspendue à un ressort dont l'autre extrémité est fixe, qui a déjà été traité au chapitre précédent.

7.1 Positions d'équilibre - Stabilité

Les deux forces, force de rappel du ressort et poids, travaillent. Il a été établi que l'énergie potentielle de la force de rappel du ressort s'exprime par :

$$E_{p_{\text{ressort}}} = \frac{1}{2} (l - l_0)^2 + C$$

en notant C une constante et que l'énergie potentielle associée au poids est :

$$E_{p_{\text{poids}}} = -mgx + C'$$

où C' est une constante, l'axe Ox étant dirigé selon la verticale descendante.

Dans le cas du mouvement horizontal, l'énergie potentielle du poids est une constante puisque l'altitude ne varie pas. En prenant comme origine des potentiels la position d'équilibre x_{eq} c'est-à-dire la position telle que la longueur du ressort soit sa longueur à vide, on obtient l'expression suivante de l'énergie potentielle totale du système :

$$E_p = \frac{1}{2} kX^2$$

où $X = x - x_{eq}$.

Dans le cas du mouvement vertical, le terme d'énergie potentielle liée au poids doit être pris en compte :

$$E_p = -mgx + \frac{1}{2} k(x - l_0) + C'$$

On choisit l'énergie potentielle nulle à la position d'équilibre $x_{eq} = l_0 + \frac{mg}{k}$. L'expression de l'énergie potentielle totale du système devient :

$$E_p = \frac{1}{2} kX^2$$

On a donc dans les deux cas la même expression et on remarque que l'énergie potentielle du poids disparaît.

On retrouve également qu'à l'équilibre, l'énergie potentielle E_p est minimale et qu'il s'agit d'une position d'équilibre stable.

7.2 Équation du mouvement

On applique le théorème de l'énergie cinétique pour retrouver l'équation du mouvement.

Dans tous les cas, la position du point matériel est repérée par rapport à la position d'équilibre.

L'énergie mécanique est :

$$E_m = E_c + E_p = \frac{1}{2}m\dot{X}^2 + \frac{1}{2}kX^2$$

soit en dérivant par rapport au temps :

$$m\dot{X}\ddot{X} + k\dot{X}X = 0$$

Comme on étudie le mouvement, l'équation $\dot{X} = 0$ est sans intérêt (point immobile).

On en déduit donc l'équation du mouvement :

$$\ddot{X} + \frac{k}{m}X = 0$$

qui est la même que celle établie par le principe fondamental de la dynamique.

A. Applications directes du cours

1. La force est-elle conservative ?

Un point $M(x, y)$ est soumis à l'action d'une force dépendant de la position selon la relation : $\vec{f} = (y^2 - x^2)\vec{u}_x + 3xy\vec{u}_y$. On donne les points $A(x_0, 0)$, $B(0, y_0)$ et $C(x_0, y_0)$.

1. Calculer le travail de $\vec{f}$ lorsque M se déplace de l'origine à C directement.
2. Même question en passant par A .
3. Même question en passant par B .
4. Conclure sur la nature conservative ou non de $\vec{f}$.

2. Travail d'une force

Un homme tire un traîneau de bas en haut d'une colline dont la forme est assimilée à un demi-cercle de rayon R , de centre O . Il exerce une force de traction $\vec{T}$ constante en norme et faisant un angle α constant avec le sol.

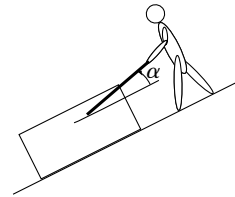


Figure 6.10

1. Déterminer le travail de la force $\vec{T}$ sur le trajet.
2. a) Sachant que l'homme se déplace à vitesse constante v , déterminer la réaction normale $\vec{R}_n$ exercée sur le traîneau en fonction de m, g, θ (l'angle polaire), T, α, R et v .
b) La loi de Coulomb donne la norme de la force de frottement du traîneau sur le sol $\|\vec{F}_f\| = f\|\vec{R}_n\|$, f étant une constante. Déterminer le travail de la force de frottement sur le trajet.

3. Le Marsupilami

Le Marsupilami est un animal de bande dessinée créé par Franquin aux capacités physiques remarquables, en particulier grâce à sa queue qui possède une force importante. Pour se déplacer, le Marsupilami enroule sa queue comme un ressort entre lui et le sol et s'en sert pour se propulser vers le haut.

On appelle $\ell_0 = 2$ m la longueur à vide du ressort équivalent. Lorsqu'il est complètement comprimé, la longueur du ressort est $\ell_m = 50$ cm. La masse m de l'animal est 50 kg et la queue quitte le sol lorsque le ressort mesure ℓ_0 . On prendra $g = 10 \text{ m.s}^{-2}$.

Quelle est la constante de raideur du ressort équivalent si la hauteur maximale d'un saut est $h = 10$ m ? Quelle est sa vitesse lorsque la queue quitte le sol ?

4. Toboggan

Un adulte ($m = 70$ kg) descend un toboggan d'une hauteur $h = 5$ m faisant un angle $\alpha = 45^\circ$ avec le sol. La norme de la force de frottement $\vec{T}$ est donnée par $\|\vec{T}\| = f\|\vec{R}\|$, où $f = 0,4$ est le coefficient de frottement et $\vec{R}$ la réaction normale. On prendra $g = 9,8 \text{ m.s}^{-2}$.

1. Calculer la variation d'énergie mécanique due au frottement entre le haut et le bas du toboggan.

2. Déterminer la vitesse de la personne en bas du toboggan. Comparer la à celle qu'il aurait s'il n'y avait pas de frottement.

5. Interaction entre particules chargées

On considère deux particules A (fixe) et B (mobile), de même masse m et de charge respective q_A et q_B . On rappelle que la force de Coulomb est $\vec{f} = \frac{q_A q_B}{4\pi\epsilon_0} \frac{\vec{AB}}{AB^3}$. C'est la seule force à prendre en compte dans cet exercice.

1. Déterminer l'énergie potentielle dont dérive la force $\vec{f}$.
2. On suppose $q_A = q_B = q$. On lance B vers A avec la vitesse v_0 . À quelle distance minimale B s'approche-t-elle de A ?
3. On suppose $q_A = -q_B = q$. Quelle vitesse minimale faut-il donner à B pour qu'elle puisse s'échapper à l'infini?

B. Exercices et problèmes

1. Cycliste au Tour de France, d'après Centrale 2006

Un cycliste assimilé à un point matériel se déplace en ligne droite. Il fournit une puissance mécanique constante P , les forces de frottement de l'air sont proportionnelles au carré de la vitesse v du cycliste ($\vec{F}_f = -kv\vec{v}$), k est une constante positive. On néglige les forces de frottement du sol sur la roue et on choisit un axe horizontal Ox .

1. En appliquant le théorème de la puissance cinétique, établir une équation différentielle en v et montrer qu'on peut la mettre sous la forme :

$$mv^2 \frac{dv}{dx} = P - kv^3$$

- a) On pose $f(x) = P - kv^3$. Dédurre des résultats précédent, l'équation différentielle vérifiée par f .
- b) Déterminer l'expression de la vitesse en fonction de x , s'il aborde la ligne droite avec une vitesse v_0 .

2. Bille dans une gouttière, d'après ICNA 2006

Une bille, assimilée à un point matériel M de masse m , est lâchée sans vitesse initiale depuis le point A d'une gouttière situé à une hauteur h du point le plus bas O de la gouttière. Cette dernière est terminée en O par un guide circulaire de rayon a , disposé verticalement. La bille, dont on suppose que le mouvement a lieu sans frottement, peut éventuellement quitter la gouttière à l'intérieur du cercle. On désigne par $\vec{g} = -g\vec{u}_y$ l'accélération de la pesanteur.

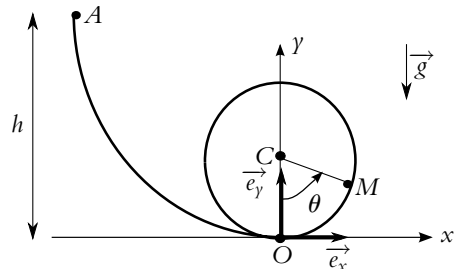


Figure 6.11

1. Calculer la norme v_0 de la vitesse en O puis en un point M quelconque du cercle repéré par l'angle θ .
2. Déterminer la réaction de la gouttière en un point du guide circulaire.
3. Déterminer la hauteur minimale de A pour que la bille ait un mouvement révolatif dans le guide.

3. Pendule

Un pendule est constitué d'un point M de masse m suspendu à un fil de longueur ℓ et de masse négligeable. Le point d'attache O du fil est pris comme origine du repère d'un axe Oz ascendant. On note v_0 la vitesse lorsque M est en $z = -\ell$.

1. Avec un raisonnement énergétique, déterminer la norme de la vitesse v lorsque M est à la cote z .
2. En déduire la tension T du fil à l'altitude z .
3. Discuter suivant la valeur de v_0 le type de mouvement obtenu : oscillatoire non interrompu (T ne s'annule pas), oscillatoire interrompu, révolatif (le pendule effectue un tour) interrompu ou non interrompu.

4. Équilibre de charges et petits mouvements

Tous les mouvements ont lieu le long de l'axe Ox et on néglige l'effet du poids. La force électrostatique exercée par un point M_1 de charge q_1 sur un point M_2 de charge q_2 est :

$$\vec{f} = \frac{q_1 q_2}{4\pi\epsilon_0} \frac{\vec{u}}{M_1 M_2^2}$$

$\vec{u}$ étant un vecteur unitaire dirigé de M_1 vers M_2 .

On travaille dans un référentiel $\mathcal{R}$ supposé galiléen.

1. Un point matériel de charge q est fixé en O , origine du repère. Un autre point matériel M de charge identique q , initialement à droite de O , à la distance a de O , est abandonné sans vitesse initiale.
 - a) Déterminer en fonction de $x = \overline{OM} > 0$ l'énergie potentielle associée à la force électrostatique subie par M . On la prendra par convention nulle pour M à l'infini.
 - b) Déterminer la vitesse de M lorsqu'il est à l'infini.
2. Un autre point matériel fixe O' de charge $4q$ est ajouté à une distance $2a$ à droite de O . Le point M est placé entre O et O' . Déterminer vectoriellement la force qui s'exerce sur M en fonction de $x = \overline{OM}$ ($0 < x < 2a$). En déduire la position d'équilibre.
3. Déterminer l'expression de l'énergie potentielle $E_p(x)$ du point et étudier la stabilité de la position d'équilibre pour $0 < x < 2a$.
4. Déterminer l'énergie mécanique du point M .
5. Représenter l'allure de $E_p(x)$ pour la particule de charge q pour $x \in [0, 2a]$. Calculer la valeur de l'extremum.

6. On lâche M sans vitesse initiale à une distance a de chacune des charges. Déterminer graphiquement puis par le calcul les distances d'approche minimales de O et O' et l'énergie cinétique maximale.

7. On lâche sans vitesse initiale le point M (de charge q) au voisinage de la position d'équilibre (C.f. question 2). Établir l'équation différentielle du mouvement au voisinage de cette position ainsi que la pulsation des petites oscillations. On rappelle que si $x \ll 1$ alors $(1+x)^\alpha \simeq 1 + \alpha x + \frac{\alpha(\alpha-1)x^2}{2}$.

5. Tige avec ressort

On considère une tige fixe, dans un plan vertical xOz , faisant l'angle α avec l'axe Oz . Un anneau M de masse m est enfilé sur la tige et astreint à se déplacer sans frottement le long de celle-ci. Cet anneau est également relié à un ressort de raideur k et de longueur à vide ℓ_v dont l'autre extrémité est fixée en O . On repère la position de M par $OM = X$.

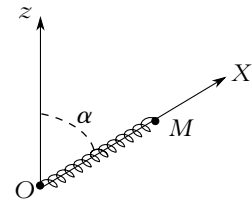


Figure 6.12

1. Quelles sont les forces conservatives appliquées à M ? Déterminer l'expression de leur énergie potentielle E_p en fonction de X et α .

2. Établir l'équation différentielle du mouvement à l'aide du théorème de l'énergie mécanique.

3. a) Étudier la fonction $E_p(X)$ dans le cas où $mg \cos(\alpha) < k\ell_v$. Tracer son allure.

b) Discuter sur le graphique les mouvements possibles en prenant à $t = 0$ les conditions initiales suivantes : $X = \ell_v$ et $\frac{dX}{dt} = V_0$. Préciser la valeur maximale de V_0 pour que le mouvement se fasse entre deux positions extrêmes $X_1 > 0$ et $X_2 > 0$.

c) Déterminer les fonctions $V(X)$ et $X(t)$ dans les conditions de la question précédente.

6. Mouvement d'une bille reliée à un ou deux ressorts sur un cercle

On considère le mouvement d'une bille M de masse m pouvant coulisser sans frottement sur un cerceau de centre O et de rayon R . On note AB le diamètre horizontal du cerceau, Ox l'axe horizontal, Oy l'axe vertical et θ l'angle entre Ox et OM . Le sens trigonométrique est pris positif.

1. Dans un premier temps, la bille est attachée à un ressort de longueur à vide nulle et de raideur k dont la seconde extrémité est fixée en B .

a) Déterminer l'angle α entre MO et MB en fonction de θ .

b) Établir l'expression de la longueur de ressort en fonction de R et θ .

c) Établir l'équation différentielle du mouvement.

d) Déterminer les positions d'équilibre.

e) On écarte la bille de sa position d'équilibre. Écrire l'équation du mouvement.

f) A quelle condition a-t-on des oscillations ?

2. Dans la suite, en plus du ressort précédent, on attache la bille à un deuxième ressort identique au premier dont la seconde extrémité est fixée en A .

- Déterminer l'angle β entre MO et MA en fonction de θ .
- Établir l'expression de la longueur MA en fonction de R et θ .
- Établir l'équation différentielle du mouvement.
- Les ressorts interviennent-ils dans le mouvement ?
- Définir la force résultant des tensions des ressorts et justifier d'une deuxième manière la réponse de la question précédente.
- Exprimer l'énergie potentielle de M en fonction de R , m , g et θ en prenant comme origine des potentiels le point le plus bas du cerceau.
- Retrouver ainsi l'équation du mouvement.
- Quelles sont les positions d'équilibre ?
- Sont-elles stables ?
- Exprimer la réaction $\vec{N}$ du cerceau en fonction de θ sachant qu'on abandonne la bille depuis $\theta = \theta_0$ sans vitesse initiale.
- Établir la condition à vérifier pour que la bille quitte le cerceau.
- En déduire une condition nécessaire en fonction de k , R , m et g pour qu'il en soit ainsi pour $\theta_0 = 0$.

7. Mouvement d'un anneau sur une piste circulaire, d'après CCP 2004

On considère le dispositif de la figure (6.13) où un objet assimilable à un point matériel M de masse m se déplace solidai-
rement à une piste formée de deux parties
circulaires de rayon R_1 et R_2 , de centre C_1
et C_2 dans un plan vertical.

On repère la position de M par l'angle θ .
Pour la partie (1), $\theta \in \left[-\frac{\pi}{2}, \pi\right]$ et pour
la partie (2), $\theta \in \left[\pi, \frac{5\pi}{2}\right]$. Il n'y a pas de
frottement et on note g l'accélération de la
pesanteur.

1. Exprimer l'énergie potentielle de
pesanteur E_p du point M en supposant
 $E_p = 0$ au point B ($\theta = \pi$). On distin-
guera les cas (1) et (2).

2. Tracer l'allure de $E_p(\theta)$.

3. Déterminer les positions angulaires
d'équilibre et leur stabilité.

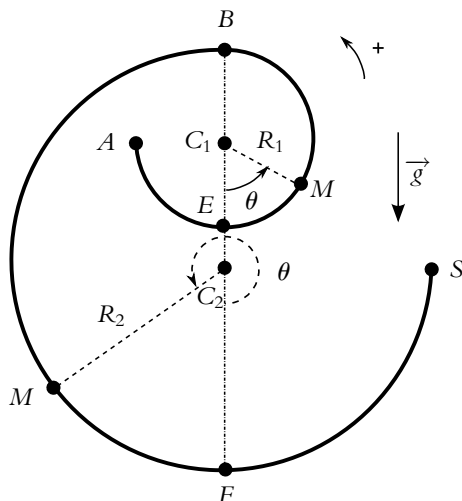


Figure 6.13

4. L'anneau est initialement en A ($\theta = -\pi/2$), il est lancé à une vitesse V_0 dans le sens trigonométrique.

a) À quelle condition sur la vitesse V_0 l'anneau peut-il atteindre le point F ($\theta = 2\pi$) ?

b) Cette condition étant remplie, donner l'expression de sa vitesse V_F en fonction des données du problème.

c) À quelle condition sur V_0 , l'anneau sort-il de la piste en S ($\theta = 5\pi/2$).

8. Pendule simple modifié, d'après ENSTIM 2005

On considère un pendule simple modifié (figure 6.14). Un mobile ponctuel M de masse m , est accroché à l'extrémité d'un fil inextensible de longueur L et de masse négligeable, dont l'autre extrémité est fixe en O . On néglige tout frottement. Lorsque $\theta > 0$, le système se comporte comme un pendule simple de centre O et de longueur de fil L . À la verticale et en dessous de O , un clou est planté en O' avec $OO' = L/3$, qui bloquera la partie haute du fil vers la gauche.

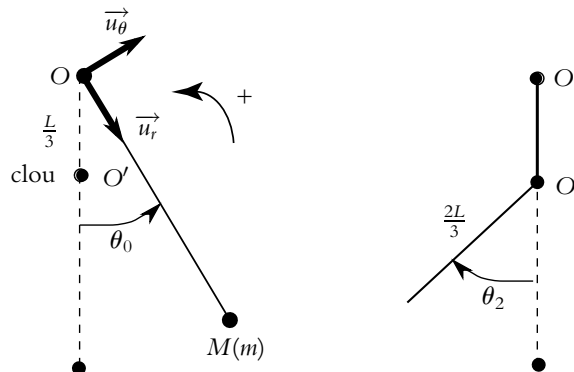


Figure 6.14

Quand $\theta < 0$, le système se comporte donc comme un pendule simple de centre O' et de longueur de fil $2L/3$. On repère la position du pendule par l'angle θ qu'il fait avec la verticale. À la date $t = 0$, on abandonne sans vitesse initiale le mobile M en donnant au fil une inclinaison initiale $\theta(0) = \theta_0 > 0$. On note t_1 la date de la première rencontre du fil avec le clou, t_2 la date de première annulation de la vitesse du mobile pour $\theta < 0$. L'intervalle de date $[0, t_1[$ est nommé première phase du mouvement, l'intervalle $[t_1, t_2]$ est nommé deuxième phase. À la date t_1^- immédiatement inférieure à t_1 , le fil n'a pas encore touché le clou et à la date t_1^+ , immédiatement supérieure, le fil vient de toucher le clou.

1. Par le théorème de la puissance mécanique, établir l'équation différentielle vérifiée par θ pour la première phase du mouvement.

2. Dans l'hypothèse des petites oscillations, on suppose $\sin \theta \simeq \theta$. Déterminer la durée δt_1 de la première phase du mouvement sans résoudre l'équation.

3. En utilisant le théorème de l'énergie mécanique, déterminer la vitesse v_1^- de M à la date t_1^- . En déduire la vitesse angulaire ω_1^- à cette date.

4. Le blocage de la partie supérieure du fil par le clou ne s'accompagne d'aucun transfert énergétique. déterminer la vitesse v_1^+ de M à la date t_1^+ . En déduire la vitesse angulaire ω_1^+ à cette date.
5. En utilisant les résultats des questions 1 et 2, donner sans calcul la durée δt_{II} de la deuxième phase.
6. Déterminer l'expression de l'angle θ_2 à la date t_2 .
7. Décrire brièvement la suite du mouvement de ce système et donner l'expression de sa période T .
8. Dresser l'allure du portrait de phase, dans le système d'axes $\left(\theta, \frac{d\theta}{dt}\right)$. Conclure.

Oscillateurs harmoniques et amortis par frottement fluide

7

L'objet de ce chapitre est l'étude des oscillateurs harmoniques et amortis par frottement fluide. On applique les résultats établis dans les chapitres précédents.

1. Oscillateurs harmoniques

Les oscillateurs harmoniques décrivent des comportements oscillants qu'ils soient dus à une nature intrinsèquement oscillatoire (comme le mouvement d'une masse reliée à un ressort) ou à un mouvement au voisinage d'une position d'équilibre (comme dans le modèle d'une liaison moléculaire). Dans les deux cas, on utilise le même modèle de l'oscillateur harmonique.

1.1 Exemples d'oscillateurs harmoniques

Aux chapitres précédents, on a étudié le mouvement d'une masse attachée à un ressort. Que le mouvement soit vertical ou horizontal, on a obtenu l'équation suivante :

$$\ddot{x} + \frac{k}{m}x = 0$$

en notant x l'allongement du ressort par rapport à sa position d'équilibre. Il s'agit de l'équation caractéristique des oscillateurs harmoniques.

Avant de décrire le mouvement correspondant, nous verrons un autre exemple conduisant au même type d'équation du mouvement : l'étude des petites oscillations au voisinage d'une position d'équilibre stable.

Considérons une particule ponctuelle de masse m soumise à la seule force $\vec{f}$ conservative. Celle-ci dérive de l'énergie potentielle Ep ne dépendant que d'un paramètre de position x . On a établi que $\vec{f} = f_x \vec{u}_x$ avec :

$$f_x = -\frac{dEp}{dx}$$

Dans le chapitre sur les aspects énergétiques, on a établi que si $Ep(x)$ admet un minimum pour $x = x_0$, il existe un point d'équilibre stable en $x = x_0$. On peut effectuer un développement limité à l'ordre 2 de l'énergie potentielle au voisinage de ce point :

$$Ep(x) = Ep(x_0) + (x - x_0) \left(\frac{dEp}{dx} \right)_{x=x_0} + \frac{(x - x_0)^2}{2} \left(\frac{d^2Ep}{dx^2} \right)_{x=x_0}$$

Le fait d'effectuer un développement limité restreint l'étude à de faibles amplitudes autour de x_0 de manière à pouvoir écrire :

$$|x - x_0| \ll |x_0|$$

Autrement dit, on se place au voisinage de l'équilibre.

Puisqu'on a un minimum en $x = x_0$,

$$\left(\frac{dEp}{dx} \right)_{x=x_0} = 0$$

et on peut assimiler la fonction $Ep(x)$ à une parabole au voisinage de l'équilibre :

$$Ep(x) = Ep(x_0) + \frac{(x - x_0)^2}{2} \left(\frac{d^2Ep}{dx^2} \right)_{x=x_0}$$

en faisant l'hypothèse que

$$\left(\frac{d^2Ep}{dx^2} \right)_{x=x_0} \neq 0$$

On notera k cette quantité.

On en déduit alors l'expression de la force au voisinage de la position d'équilibre :

$$f_x = -\frac{dEp}{dx} = -k(x - x_0)$$

On peut d'ores et déjà noter l'analogie évidente entre l'expression de cette force et celle de la force de rappel d'un ressort.

La projection du principe fondamental sur l'axe Ox donne :

$$m\ddot{x} = -k(x - x_0)$$

soit

$$\ddot{u} + \frac{k}{m}u = 0$$

en posant $u = x - x_0$. On retrouve la même équation différentielle que dans les deux exemples précédents. Il s'agit de celle d'un oscillateur harmonique.

1.2 Étude du mouvement d'un oscillateur harmonique

On note l'équation différentielle sous la forme :

$$\ddot{x} + \omega_0^2 x = 0$$

quel que soit le système dont elle est issue. Dans les différents exemples envisagés précédemment, la constante étant toujours positive, on la prend sous la forme d'un réel au carré sans que cela n'apporte de conditions supplémentaires. On peut noter que dans le cas d'un point matériel relié à un ressort, $\omega_0 = \sqrt{\frac{k}{m}}$.

a) Caractéristiques du mouvement

On rappelle que la solution de cette équation différentielle s'écrit :

$$x(t) = A \cos \omega_0 t + B \sin \omega_0 t$$

où A et B sont des constantes réelles à déterminer à partir des conditions initiales. On peut également la mettre sous la forme :

$$x(t) = X \cos(\omega_0 t + \varphi)$$

avec $X = \sqrt{A^2 + B^2}$, $\cos \varphi = \frac{A}{X}$ et $\sin \varphi = -\frac{B}{X}$.

- X désigne l'amplitude du mouvement,
- ω_0 sa pulsation,
- φ la phase initiale.

A et B , ou X et φ , sont des constantes à déterminer à partir des conditions initiales. On peut noter qu'il faut deux conditions initiales pour obtenir les deux constantes, ce pourra être la position et la vitesse initiales.

Le mouvement est périodique, de période $T = \frac{2\pi}{\omega_0}$ indépendante de l'amplitude du mouvement : on dit qu'il y a **isochronisme des oscillations**.

b) Aspect énergétique

Dans ce paragraphe, on utilise l'expression de la solution sous la forme :

$$x(t) = X \cos(\omega t + \varphi).$$

La seule force considérée ici est une force conservative dérivant du potentiel

$$E_p = \frac{1}{2} k x^2$$

E_p étant défini à une constante près, on a choisi $E_p(0) = 0$. L'énergie potentielle du système s'écrit donc :

$$E_p = \frac{1}{2}kx^2 = \frac{1}{2}kX^2 \cos^2(\omega_0 t + \varphi)$$

L'énergie cinétique est par définition :

$$E_c = \frac{1}{2}m\dot{x}^2$$

Or $\dot{x} = -X\omega_0 \sin(\omega_0 t + \varphi)$ donc

$$E_c = \frac{1}{2}mX^2\omega_0^2 \sin^2(\omega_0 t + \varphi)$$

Comme $\omega_0^2 = \frac{k}{m}$, on peut écrire :

$$E_c = \frac{1}{2}kX^2 \sin^2(\omega_0 t + \varphi)$$

L'énergie mécanique, somme des énergies potentielle et cinétique, s'écrit donc :

$$\begin{aligned} E_m &= E_p + E_c = \frac{1}{2}kx^2 + \frac{1}{2}m\dot{x}^2 \\ &= \frac{1}{2}kX^2 \cos^2(\omega_0 t + \varphi) + \frac{1}{2}kX^2 \sin^2(\omega_0 t + \varphi) = \frac{1}{2}kX^2 \end{aligned}$$

L'énergie mécanique est une constante du mouvement, ce qui n'a rien de surprenant puisque toutes les forces sont conservatives.

Les valeurs moyennes des énergies potentielle et cinétique sont :

$$\langle E_p \rangle = \frac{1}{T} \int_0^T E_p(t) dt = \frac{1}{T} \int_0^T \frac{1}{2}kX^2 \cos^2(\omega_0 t + \varphi) dt = \frac{1}{4}kX^2$$

et

$$\langle E_c \rangle = \frac{1}{T} \int_0^T E_c(t) dt = \frac{1}{T} \int_0^T \frac{1}{2}kX^2 \sin^2(\omega_0 t + \varphi) dt = \frac{1}{4}kX^2$$

Les valeurs moyennes des énergies potentielle et cinétique sont égales, il y a en permanence un transfert d'énergie entre les deux formes, potentielle et cinétique.

1.3 Portrait de phase de l'oscillateur harmonique

L'équation du mouvement $m\ddot{x} + kx = 0$ peut se mettre sous la forme :

$$\begin{cases} \dot{v} = g(x, \dot{x}) = -\frac{k}{m}x \\ \dot{x} = v \end{cases}$$

On peut donc chercher à établir l'allure du portrait de phase.

Pour cela, on dispose de deux méthodes :

- Utilisation de l'intégrale première du mouvement :

On utilise la conservation de l'énergie mécanique et on obtient :

$$\frac{1}{2}v^2 + \frac{1}{2}\omega_0^2 x^2 = C$$

où C désigne une constante qui varie avec les conditions initiales. En notant $A = 2C > 0$ et $B = \frac{2C}{\omega_0^2} > 0$, l'équation de la courbe se met sous la forme :

$$\frac{v^2}{A} + \frac{x^2}{B} = 1$$

qui est l'équation cartésienne d'une ellipse.

- Explicitation de la solution :

La résolution de l'équation différentielle conduit à la solution :

$$x = A \cos(\omega_0 t + \varphi)$$

où A et φ sont des constantes variant avec les conditions initiales¹. On en déduit $v = \dot{x} = -A\omega_0 \sin(\omega_0 t + \varphi)$ et comme $\sin^2 x + \cos^2 x = 1$ pour tout x on a :

$$\frac{v^2}{A^2 \omega_0^2} + \frac{x^2}{A^2} = 1$$

On retrouve l'équation cartésienne d'une ellipse.

Lorsque la constante C est nulle, la trajectoire de phase est réduite au point de coordonnées $x = 0$ et $v = 0$. C'est le centre des différentes ellipses, il correspond à la position d'équilibre du système.

Le portrait de phase d'un oscillateur harmonique est donc un ensemble d'ellipses de centre commun qui se déduisent les unes des autres par homothétie.

On peut remarquer que :

- en traçant $(x, \frac{v}{\omega_0})$, les trajectoires de phase sont des cercles,
- dans le plan (x, v) , les trajectoires de phase seront des cercles lorsque $\omega_0 = 1 \text{ rad.s}^{-1}$, ce qui correspond à une fréquence $f_0 = \frac{\omega_0}{2\pi} = 0,159 \text{ Hz}$.

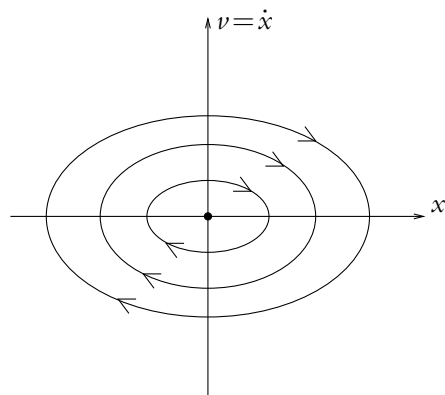


Figure 7.1 Portrait de phase de l'oscillateur harmonique.

¹ Si les conditions initiales sont x_0 et v_0 alors $\varphi = -\text{Arctan} \frac{v_0}{\omega_0 x_0}$ et $A = \frac{x_0}{\cos \varphi}$.

2. Oscillateurs amortis

Dans le paragraphe précédent, on n'a pas tenu compte d'éventuelles forces de frottement. Or les frottements existent toujours mais se font plus ou moins ressentir. On s'intéresse ici au cas où on a un frottement fluide proportionnel à la vitesse dont la force s'écrit :

$$\vec{f} = -\lambda \vec{v}$$

2.1 Équation différentielle du mouvement

On établit l'équation différentielle du mouvement sur l'exemple d'un oscillateur à ressort.

On procède toujours avec les quatre étapes :

1. Système : point matériel de masse m se déplaçant le long d'un axe Ox , dont l'origine de l'axe est fixée au niveau de la position à l'équilibre du point matériel,
2. Référentiel du laboratoire considéré comme galiléen,
3. Bilan des forces :
 - la force de rappel $-kx\vec{u}_x$,
 - la force de frottement $-\lambda\dot{x}\vec{u}_x$,
 - le poids qui est compensé soit par la réaction de l'axe dans le cas où le pendule est horizontal soit « par la position d'équilibre » dans le cas où le pendule est vertical.
4. Résolution par application du principe fondamental de la dynamique, projeté sur l'axe Ox :

$$m\ddot{x} = -kx - \lambda\dot{x}$$

L'équation différentielle du mouvement est donc

$$\ddot{x} + \frac{\lambda}{m}\dot{x} + \frac{k}{m}x = 0$$

2.2 Analogie avec le circuit électrique R, L, C série

On a déjà rencontré ce type d'équation différentielle dans le cours d'électricité lors de l'étude du circuit R, L, C série en court-circuit. Les phénomènes d'oscillations sont donc présents dans divers domaines de la physique comme la mécanique et l'électrocinétique.

Pour donner un sens physique à l'analogie qui va être établie, il faut réécrire l'équation différentielle électrique en fonction de la charge q de la capacité au lieu de sa tension u et considérer le circuit R, L, C série en court-circuit. On a la relation $q = Cu$, ce qui permet d'obtenir :

$$L\ddot{q} + R\dot{q} + \frac{q}{C} = 0$$

Le tableau suivant explicite l'analogie qui peut être faite entre la mécanique et l'électrocinétique :

	Mécanique	Électronique
Équation différentielle	$m\ddot{x} + \lambda\dot{x} + kx = 0$	$L\ddot{q} + R\dot{q} + \frac{1}{C}q = 0$
Pulsation propre	$\sqrt{\frac{k}{m}}$	$\sqrt{\frac{1}{LC}}$
Variables	élongation x vitesse $v = \dot{x}$ masse m ressort de raideur k frottement proportionnel à v de coefficient λ énergie cinétique $\frac{1}{2}m\dot{x}^2$ énergie potentielle $\frac{1}{2}kx^2$	charge q intensité i inductance L capacité $\frac{1}{C}$ résistance R énergie magnétique $\frac{1}{2}L\dot{q}^2$ énergie électrostatique $\frac{1}{2}\frac{1}{C}q^2$

Dès qu'on aura, quel que soit le domaine concerné, une équation différentielle du type oscillateur amorti, l'ensemble des résultats pourra être obtenu par simple analogie comme celle qu'on vient d'établir entre le pendule mécanique amorti et le circuit R, L, C série en court-circuit. L'étude que l'on va refaire en mécanique n'est pas nécessaire : il suffit d'établir les résultats par analogie avec ce qui a été fait en électrocinétique.

2.3 Trois types de régimes

a) Équation caractéristique

L'équation du mouvement s'écrit :

$$\ddot{x} + \frac{\lambda}{m}\dot{x} + \frac{k}{m}x = 0$$

On a l'habitude de mettre cette équation sous une forme dite canonique :

$$\ddot{x} + 2\xi\dot{x} + \omega_0^2x = 0 \tag{1}$$

Ici cela revient à poser $\omega_0 = \sqrt{\frac{k}{m}}$ et $\xi = \frac{\lambda}{2m}$. On utilisera cette notation dans toute la suite. On notera qu'en électricité, on a utilisé une définition différente pour ξ : le

coefficient de la dérivée première était notée $2\xi\omega_0$ au lieu de 2ξ . Ici cela reviendrait à poser $\xi = \frac{\lambda}{\sqrt{km}}$. Il faudra effectuer ce remplacement pour retrouver rigoureusement les mêmes expressions.

D'autres écritures pourraient être adoptées :

- en fonction de ω_0 et d'un temps de relaxation τ : on introduit la quantité τ

$$\tau = \frac{m}{\lambda} = \frac{1}{2\xi}$$

qui est homogène à un temps et qu'on appelle temps de relaxation. L'équation différentielle peut alors se réécrire :

$$\ddot{x} + \frac{1}{\tau}\dot{x} + \omega_0^2 x = 0$$

- en fonction de ω_0 et du facteur de qualité Q : le facteur de qualité est défini comme la quantité sans dimension qu'il est possible de former à l'aide de la pulsation propre et du temps de relaxation :

$$Q = \omega_0 \tau = 2\pi \frac{\tau}{T_0}$$

On peut alors écrire l'équation différentielle sous la forme :

$$\ddot{x} + \frac{\omega_0}{Q}\dot{x} + \omega_0^2 x = 0$$

Pour résoudre l'équation différentielle sous la forme (1), on introduit l'équation caractéristique associée :

$$r^2 + 2\xi r + \omega_0^2 = 0$$

Si on appelle r_1 et r_2 les solutions de l'équation caractéristique, la solution s'écrit :

$$x(t) = Ae^{r_1 t} + Be^{r_2 t}$$

Plusieurs types de solutions sont distingués suivant le signe du discriminant réduit de l'équation caractéristique :

$$\Delta' = \xi^2 - \omega_0^2$$

b) Régime pseudopériodique

Il s'agit du cas où le discriminant est négatif :

$$\xi < \omega_0 \Leftrightarrow \frac{\lambda}{2m} < \sqrt{\frac{k}{m}} \Leftrightarrow \omega_0 \tau > \frac{1}{2} \Leftrightarrow Q > \frac{1}{2}$$

Cela correspond donc au cas où le frottement défini grâce au coefficient λ est faible. On peut noter que l'oscillateur harmonique correspond à la limite $\lambda = 0$.

Les deux racines de l'équation caractéristique sont complexes conjuguées :

$$r_{\pm} = -\xi \pm i\sqrt{\omega_0^2 - \xi^2} = -\xi \pm i\omega$$

La solution de l'équation différentielle s'écrit donc :

$$x = \exp(-\xi t) (X_1 \cos \omega t + X_2 \sin \omega t) = X \exp(-\xi t) \cos(\omega t + \varphi)$$

où (X_1, X_2) ou (X, φ) sont des couples de constantes à déterminer à partir des conditions initiales.

On peut écrire la solution sous la forme du produit d'une amplitude dépendant du temps $A(t) = X \exp(-\xi t)$ et d'un terme d'oscillations $\cos(\omega t + \varphi)$. Des oscillations sont observées mais leur amplitude $A(t)$ diminue exponentiellement au fur et à mesure qu'ont lieu les oscillations. On verra que cette décroissance correspond à une dissipation d'énergie due au frottement.

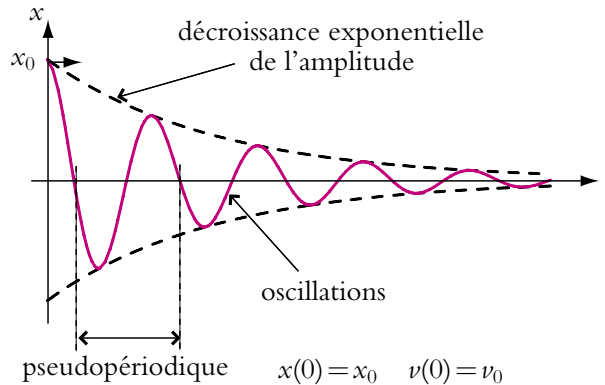


Figure 7.2 Régime pseudopériodique de l'oscillateur amorti.

Pour ce type de mouvement, on définit deux grandeurs : la pseudopériode et le décrément logarithmique.

Pseudopériode

C'est l'intervalle de temps entre trois passages par la position d'équilibre, à savoir la période du $\cos(\omega t + \varphi)$ qui vaut :

$$T = \frac{2\pi}{\omega} = \frac{2\pi}{\omega_0 \sqrt{1 - \frac{\xi^2}{\omega_0^2}}}$$

Or, dans le cas étudié ici, $\frac{\xi}{\omega_0} = \frac{\lambda}{2m\omega_0} < 1$, on en déduit que $T > T_0$ en notant

$T_0 = \frac{2\pi}{\omega_0}$ la période de l'oscillateur harmonique associé ou période propre.

Dans le cas d'un amortissement très faible $\frac{\xi}{\omega_0} \ll 1$, on peut effectuer un développement limité de la pseudopériode à l'ordre 2 en $\frac{\alpha}{\omega_0}$:

$$T \simeq \frac{2\pi}{\omega_0} \left(1 + \frac{\xi^2}{2\omega_0^2} + \dots \right) \simeq T_0$$

On peut identifier la pseudopériode à la période propre dans le cas d'un amortissement très faible.

À la limite où il n'y a pas de frottement, on retrouve l'égalité avec la période, ce qui est normal et rassurant.

Décrément logarithmique

Il caractérise l'amortissement :

$$\delta = \ln \frac{A(t)}{A(t+T)} = \ln \frac{A(t)}{A(t) \exp(-\xi T)} = \xi T = \frac{2\pi}{\omega_0 \sqrt{1 - \frac{\xi^2}{\omega_0^2}}}$$

c) Régime apériodique

Il correspond au cas $\Delta' > 0$ donc à un fort frottement. L'équation caractéristique admet alors deux racines réelles :

$$r_{\pm} = -\xi \pm \sqrt{\xi^2 - \omega_0^2}$$

Comme $\xi < \sqrt{\xi^2 - \omega_0^2}$, les deux solutions sont négatives.

La solution de l'équation différentielle :

$$x = X_1 \exp(r_+ t) + X_2 \exp(r_- t)$$

est la somme de deux exponentielles décroissantes. L'amplitude diminue au cours du temps sans que n'apparaissent d'oscillations, d'où le nom donné de *régime apériodique* à ce type de solutions.

d) Régime critique

Ce régime correspond au cas $\Delta = 0$. L'équation caractéristique admet

$$r = -\xi = -\omega_0$$

comme racine double. La solution s'écrit alors

$$x = (X_1 t + X_2) \exp(-\omega_0 t)$$

où X_1 et X_2 sont des constantes à déterminer à partir des conditions initiales.

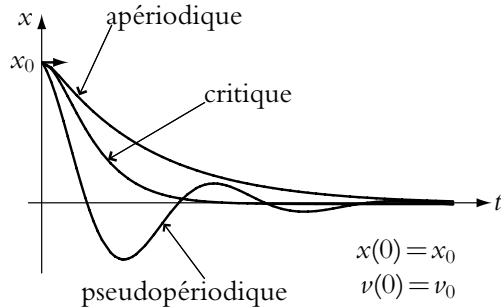


Figure 7.3 Comparaison des régimes d'un oscillateur amorti.

L'allure est la même que pour le régime apériodique. Dans ce régime, on n'observe pas d'oscillations et le retour à la position d'équilibre est plus rapide que dans le cas d'un régime pseudopériodique. D'autre part, il s'agit d'un régime théorique : d'un point de vue physique, il est la limite entre les deux autres cas et ne sera jamais observé car l'égalité ne sera jamais rigoureusement obtenue.

2.4 Aspects énergétiques

On limite l'étude au cas où on a effectivement des oscillations donc au régime pseudopériodique :

$$x = X \exp(-\xi t) \cos(\omega t + \varphi)$$

avec $\xi = \frac{\lambda}{2m}$, $\omega_0^2 = \frac{k}{m}$ et $\omega = \omega_0 \sqrt{1 - \frac{\xi^2}{\omega_0^2}}$.

La vitesse a pour expression :

$$\dot{x} = X \exp(-\xi t) \left(-\xi \cos(\omega t + \varphi) - \omega \sin(\omega t + \varphi) \right)$$

a) Expression des énergies cinétique, potentielle et mécanique

Les différentes énergies sont :

- énergie cinétique :

$$E_c = \frac{1}{2} m \dot{x}^2 = \frac{1}{2} m X^2 \exp(-2\xi t) \left(\xi \cos(\omega t + \varphi) + \omega \sin(\omega t + \varphi) \right)^2$$

- énergie potentielle :

$$E_p = \frac{1}{2} k x^2 = \frac{1}{2} k X^2 \exp(-2\xi t) \cos^2(\omega t + \varphi) = \frac{1}{2} m \omega_0^2 X^2 \exp(-2\xi t) \cos^2(\omega t + \varphi)$$

car $k = m \omega_0^2$

- énergie mécanique :

$$\begin{aligned} E_m &= E_c + E_p \\ &= \frac{1}{2} m X^2 \exp(-2\xi t) \left((\xi \cos(\omega t + \varphi) + \omega \sin(\omega t + \varphi))^2 + \omega_0^2 \cos^2(\omega t + \varphi) \right) \end{aligned}$$

L'énergie mécanique n'est donc pas constante. Elle diminue en moyenne selon une loi en $\exp(-2\xi t)$: il y a dissipation d'énergie.

b) Cas d'un amortissement très faible

On se place dans le cas où

$$\xi \ll \omega_0$$

On a déjà établi qu'alors on pouvait assimiler pseudo-pulsation et pulsation propre :

$$\omega \simeq \omega_0$$

De cette hypothèse sont déduites les expressions approchées de :

- l'énergie cinétique :

$$Ec \simeq \frac{1}{2} m \omega_0^2 X^2 \exp(-2\xi t) \omega_0^2 \sin^2(\omega_0 t + \varphi)$$

- l'énergie potentielle :

$$Ep \simeq \frac{1}{2} k X^2 \exp(-2\xi t) \cos^2(\omega_0 t + \varphi) = \frac{1}{2} m \omega_0^2 X^2 \exp(-2\xi t) \cos^2(\omega_0 t + \varphi)$$

- l'énergie mécanique :

$$Em = \frac{1}{2} m \omega_0^2 X^2 \exp(-2\xi t) = \frac{1}{2} m \omega_0^2 X^2 \exp\left(-\frac{t}{\tau}\right)$$

On peut alors estimer la perte d'énergie au cours d'une pseudopériode assimilée ici à la période propre. On remarque que $Em(t + T_0) < Em(t)$. Pour avoir une variation d'énergie positive, on considère donc $\Delta Em = Em(t) - Em(t + T_0)$:

$$\Delta Em = Em(t) - Em(t + T_0) = \frac{1}{2} m \omega_0^2 X^2 \exp\left(-\frac{t}{\tau}\right) \left(1 - \exp\left(-\frac{T_0}{\tau}\right)\right) > 0$$

La variation relative d'énergie mécanique s'écrit :

$$\frac{\Delta Em}{Em} = \frac{\frac{1}{2} m \omega_0^2 X^2 \exp\left(-\frac{t}{\tau}\right) \left(1 - \exp\left(-\frac{T_0}{\tau}\right)\right)}{\frac{1}{2} m \omega_0^2 X^2 \exp\left(-\frac{t}{\tau}\right)} = \left(1 - \exp\left(-\frac{T_0}{\tau}\right)\right)$$

Or dans le cas d'un amortissement très faible :

$$\frac{T}{\tau} = \frac{2\pi \sqrt{\frac{m}{k}}}{\frac{m}{\lambda}} = \frac{2\pi \lambda}{m \omega_0} \ll 1$$

ce qui permet d'effectuer un développement limité et d'écrire au premier ordre :

$$\frac{\Delta Em}{Em} = \frac{T}{\tau} = \frac{2\pi}{\omega_0 \tau} = \frac{2\pi}{Q}$$

On obtient donc une interprétation énergétique du facteur de qualité dans le cas d'un amortissement très faible :

$$Q = 2\pi \frac{Em}{\Delta Em}$$

Cette expression n'est valable que pour un très faible amortissement donc quand $\xi \ll \omega_0$.

c) Origine de la variation d'énergie mécanique

En appliquant le théorème de l'énergie cinétique, on obtient :

$$\Delta(E_c + E_p) = W' = - \int \lambda \dot{x}^2 dt < 0$$

L'énergie mécanique du point matériel diminue : l'énergie est perdue à cause des frottements et est dissipée sous forme de chaleur.

2.5 Portrait de phase de l'oscillateur amorti

Il faut distinguer les différents types de mouvement en fonction de l'intensité de l'amortissement. On notera l'équation différentielle sous la forme :

$$\ddot{x} + \frac{\omega_0}{Q} \dot{x} + \omega_0^2 x = 0$$

Le discriminant de l'équation caractéristique s'écrit alors :

$$\Delta = \omega_0^2 \left(\frac{1}{Q^2} - 4 \right)$$

a) Cas du régime pseudopériodique

Il s'agit du cas où $\Delta < 0$ soit $Q > \frac{1}{2}$. Physiquement cela correspond à un amortissement faible. La solution s'écrit alors

$$x = X \exp\left(-\frac{\omega_0}{Q}t\right) \cos(\omega t + \varphi)$$

avec $\omega = \omega_0 \sqrt{1 - \frac{1}{4Q^2}}$. On en déduit :

$$v = \dot{x} = X \exp\left(-\frac{\omega_0}{Q}t\right) \left(-\frac{\omega_0}{Q} \cos(\omega t + \varphi) - \omega \sin(\omega t + \varphi)\right)$$

On étudie donc la courbe paramétrée :

$$\begin{cases} x = X \exp\left(-\frac{\omega_0}{Q}t\right) \cos(\omega t + \varphi) \\ v = X \exp\left(-\frac{\omega_0}{Q}t\right) \left(-\frac{\omega_0}{Q} \cos(\omega t + \varphi) - \omega \sin(\omega t + \varphi)\right) \end{cases}$$

Étudions temporairement la courbe $\dot{y}(y)$ où :

$$y = \frac{x}{X \exp\left(-\frac{\omega_0}{Q}t\right)} \quad \text{et} \quad \dot{y} = \frac{\dot{x}}{X \exp\left(-\frac{\omega_0}{Q}t\right)}$$

En effet, on se ramène ainsi au portrait de phase d'un oscillateur non amorti. On a alors établi qu'une trajectoire de phase était une ellipse : la courbe $\dot{y}(y)$ est donc une ellipse.

Le portrait de phase de l'oscillateur amorti est donc une ellipse dont la taille diminue exponentiellement au cours du temps. On observe donc une spirale.

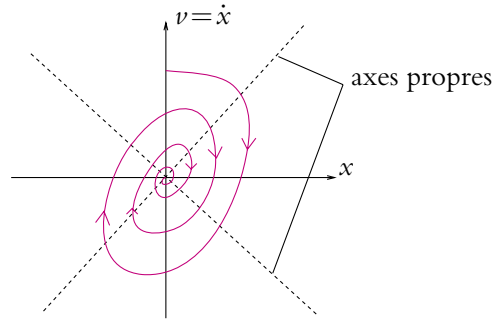


Figure 7.4 Trajectoire de phase du régime pseudopériodique.

► **Remarque** : Si l'amortissement est très faible alors $Q \gg 1$, $\omega \simeq \omega_0$ et

$$\begin{cases} x = X \exp\left(-\frac{\omega_0}{Q}t\right) \cos(\omega t + \varphi) \\ \dot{x} \simeq -\omega X \exp\left(-\frac{\omega_0}{Q}t\right) \sin(\omega t + \varphi) \end{cases}$$

b) Cas du régime critique

La solution s'écrit ici :

$$\begin{cases} x = (X_1 t + X_2) \exp(-\omega_0 t) \\ \dot{x} = (X_1 - \omega_0 X_2 - \omega_0 X_1 t) \exp(-\omega_0 t) \end{cases}$$

On note que $X_2 = x_0$ et $X_1 = v_0 + \omega_0 x_0$.

Suivant les différentes valeurs possibles de conditions initiales, l'allure de la trajectoire de phase est légèrement différente. Pour la tracer, il suffit de regarder conjointement l'évolution de x et de sa dérivée au cours du temps. C'est ce qui est fait sur le schéma suivant pour un jeu de conditions initiales tel que $X_2 > 0$ et $X_1 > \omega_0 X_2$:

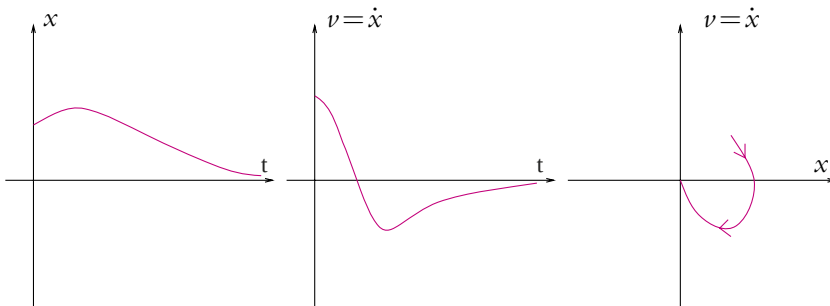


Figure 7.5 Trajectoire de phase du régime critique.

Suivant les valeurs des conditions initiales, on peut observer l'une des trajectoires de phase ci-contre.

On constate que la trajectoire de phase tend très rapidement vers la position d'équilibre à l'origine des axes et ne fait pas plusieurs rotations autour de l'origine conformément à l'absence d'oscillations dans ce régime.

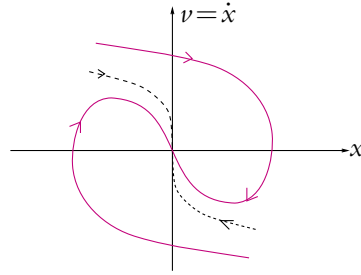


Figure 7.6 Portrait de phase du régime critique.

c) Cas du régime apériodique

On peut procéder de la même façon que pour le régime critique ou faire l'analyse suivante.

La solution s'écrit ici :

$$\begin{cases} x = X_1 \exp(\omega_1 t) + X_2 \exp(\omega_2 t) \\ v = \omega_1 X_1 \exp(\omega_1 t) + \omega_2 X_2 \exp(\omega_2 t) \end{cases}$$

$$\text{avec } \omega_1 = -\frac{\omega_0}{Q} - \omega_0 \sqrt{\frac{1}{Q^2} - 4} \text{ et } \omega_2 = -\frac{\omega_0}{Q} + \omega_0 \sqrt{\frac{1}{Q^2} - 4}.$$

On peut noter que x et v ne peuvent s'annuler que si X_1 et X_2 sont de signes contraires et que ce phénomène ne se produit qu'à un instant précis, distinct pour chaque variable.

De plus, les deux variables tendent vers 0 quand le temps devient infini : les trajectoires de phase convergent vers l'origine.

Comme ω_1 et ω_2 sont négatifs, x et v ne peuvent pas être de même signe à l'instant initial.

De ces éléments, on peut déduire l'allure du portrait de phase ci-contre.

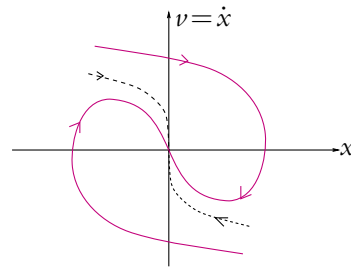


Figure 7.7 Portrait de phase du régime apériodique.

d) Amortissement ou amplification

Dans ce qui précède, l'origine du plan de phases correspond à une position d'équilibre. Celle-ci est stable : toutes les trajectoires, quel que soit le régime considéré, convergent vers l'origine en « s'enroulant » plus ou moins autour de ce point.

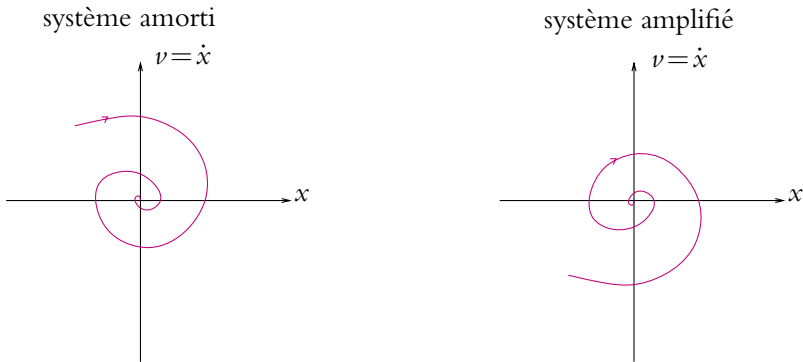


Figure 7.8 Amortissement ou amplification.

On peut noter que cette convergence vers l'origine provient du signe positif du coefficient relatif à $\dot{x}$. Il est donc possible d'imaginer des cas où ce coefficient puisse devenir négatif. Lorsque ce sera le cas, les exponentielles ne seront pas décroissantes mais croissantes. L'origine deviendra alors un point d'équilibre instable autour duquel les trajectoires de phase se « dérouleront ». On n'aura plus amortissement mais amplification.

On retrouve sur cet exemple les résultats généraux sur l'influence du frottement et sur l'amplification lors de l'étude des portraits de phase.

A. Applications directes du cours

1. Ressorts en parallèle et en série

1. On assimile la double roue de certains véhicules à deux ressorts semblables en parallèle (figure 7.9) de constante de raideur k et de longueur à vide l_0 . Déterminer les caractéristiques (k_e, l_{e0}) du ressort équivalent. On note m la masse du véhicule appuyant sur l'ensemble des deux ressorts.
2. Maintenant, si l'on s'intéresse seulement à une roue et son amortisseur (sans tenir compte du système d'amortissement visqueux), on peut assimiler l'ensemble à deux ressorts (k_1, l_{01}) et (k_2, l_{02}) en série (figure 7.10). Déterminer les caractéristiques (k_e, l_{e0}) du ressort équivalent. On note m la masse du véhicule appuyant sur l'ensemble des deux ressorts et A le point de masse nulle entre les deux ressorts.

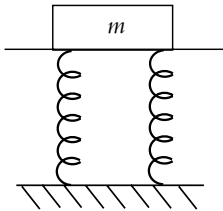


Figure 7.9

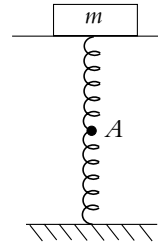


Figure 7.10

2. Ressort soumis à une force constante

Un point M de masse m est accrochée à un ressort horizontal de raideur k et de longueur à vide ℓ_0 . Le point M est astreint à un déplacement horizontal sur l'axe Ox . Un dispositif de freinage exerce une force de frottement visqueux $\vec{F}_f = -\lambda \vec{v}$ où $\vec{v}$ est la vitesse de M . À partir du repos ($x = \ell_0$), M est soumis à une force constante $\vec{F}$.

1. Établir l'équation différentielle du mouvement.
2. Quelle est la condition pour avoir une solution pseudo-périodique? Déterminer alors la solution $x(t)$.
3. On suppose que l'amortissement a été réglé pour que M n'effectue qu'une oscillation avant de s'immobiliser. Déterminer x au bout d'une pseudo-période.

3. Système de deux point couplés par des ressorts

Deux points matériels A et B de même masse m sont reliés entre eux par un ressort de raideur K , et à deux points fixes par deux ressorts de raideur k . L'ensemble coulisse sans frottement sur une tige horizontale fixe.

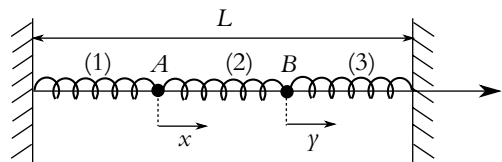


Figure 7.11

On note x et y les élongations de A et B comptées à partir de leur position d'équilibre.

1. Écrire les relations à l'équilibre reliant les longueurs à vide des ressorts (ℓ_0 pour (1) et (3) L_0 pour (2)) et les longueurs à l'équilibre (ℓ_e pour (1) et (3) L_e pour (2)).
2. Déterminer les équations du mouvement. Dédire des résultats précédents deux équations différentielles, fonction de x et y et de leurs dérivées secondes, et de k , K et m uniquement.
3. On pose $S = x + y$ et $D = x - y$. En déduire deux équations différentielles, l'une uniquement en S et l'autre uniquement en D . On posera $\omega_0^2 = k/m$ et $\omega_1^2 = (k + 2K)/m$. Déterminer les expressions générales de $S(t)$ et $D(t)$.
4. Déterminer $x(t)$ et $y(t)$ sachant qu'à $t = 0$, on écarte B de b de sa position d'équilibre, A étant maintenu à sa position d'équilibre, sans leur communiquer de vitesse initiale.

4. Oscillateur harmonique spatial

On considère un point matériel M de masse m soumis à une force $\vec{f}$ conservative d'énergie potentielle $E_p = \frac{1}{2} (k_x x^2 + k_y y^2 + k_z z^2)$.

1. Donner l'expression de la force $\vec{f}$.
2. Établir l'équation du mouvement et en déduire les équations horaires.
3. Montrer que l'énergie mécanique peut s'écrire sous la forme de la somme des énergies mécaniques d'oscillateurs harmoniques indépendants.
4. On suppose dans la suite que l'oscillateur est isotrope c'est-à-dire que $k_x = k_y = k_z$. Déterminer dans ce cas l'expression vectorielle du vecteur position $\vec{OM}$. On notera $\vec{v}_0$ la vitesse initiale du point M et $\vec{r}_0$ sa position initiale.
5. Montrer que le mouvement est plan.
6. Dans ce plan, établir que l'équation du mouvement est une ellipse.

On rappelle que l'équation cartésienne d'une ellipse peut s'écrire : $\frac{x^2}{A^2} + \frac{y^2}{B^2} + \frac{xy}{C} = 1$ avec A , B et C trois constantes positives ou négatives.

B. Exercices et problèmes

1. Oscillations d'un pendule simple, d'après CCP TSI 2005

Un objet ponctuel A de masse m est suspendu à l'extrémité d'un fil de masse négligeable et de longueur L dont l'autre extrémité O est fixe. On ne considère pas les mouvements en dehors d'un plan vertical Oxy perpendiculaire à un axe Oz horizontal. On repère A par l'angle θ entre le fil et la verticale. On suppose que le référentiel terrestre est galiléen. On néglige les frottements autour de l'axe de rotation. On désigne par $\vec{g} = g\vec{u}_x$ l'accélération du champ de pesanteur.

1. Effectuer le bilan des forces s'exerçant sur A .

2. Déterminer l'équation différentielle du mouvement de A en appliquant le principe fondamental de la dynamique.
3. Retrouver le résultat par le théorème de l'énergie cinétique.
4. Déterminer les positions d'équilibre du système.
5. Étudier leur stabilité.
6. Initialement on abandonne A sans vitesse initiale alors que le fil est écarté d'un angle θ_0 . On se place dans le cas où θ_0 est petit. Montrer que le système constitue un oscillateur harmonique dont on exprimera la pulsation ω_0 et la période T_0 en fonction de g et L .
7. Compte tenu des conditions initiales, déterminer l'expression θ de l'angle en fonction du temps t .
8. Quelle est la valeur maximale $v_{\max}$ de la vitesse de l'objet A au cours de son mouvement ? On l'exprimera en fonction de θ_0 , L et g .
9. Tracer la courbe donnant θ en fonction du temps.
10. On suppose désormais que A est soumis à une force de frottement fluide proportionnel à la vitesse $\vec{v}$. On note h le coefficient de proportionnalité.
Établir l'équation différentielle du mouvement en θ à partir du principe fondamental de la dynamique.
11. Les frottements sont suffisamment faibles pour que le régime soit pseudo-périodique. Donner la condition sur h pour qu'il en soit ainsi.
12. Dans le cas des petites oscillations, déterminer l'expression de θ en fonction du temps t . Les conditions initiales sont les mêmes que précédemment.
13. Donner l'allure de la courbe représentant θ en fonction du temps.
14. Retrouver l'équation différentielle du mouvement par une méthode énergétique.

2. Période des oscillations d'un pendule simple en fonction de l'amplitude

On considère une masse ponctuelle suspendue au bout d'un fil inextensible de masse négligeable et de longueur l .

1. Établir l'équation du mouvement de la masse m .
2. Que peut-on dire de la nature du mouvement si on ne considère que les petites oscillations ?
3. Pour de plus fortes amplitudes, on va tenir compte d'un terme supplémentaire dans le développement limité du sinus. Que devient l'équation du mouvement ?
4. On cherche des solutions de la forme : $\theta = A \cos \omega t$. Établir qu'il convient également d'introduire un terme de pulsation 3ω si on veut obtenir un mouvement.
5. On pose donc $\theta = A (\cos \omega t + \epsilon \cos 3\omega t)$ avec $A \ll 1$ et $\epsilon \ll 1$. Déterminer le système de deux équations en ω^2 et ϵ .
6. En déduire la formule de Borda $T \simeq T_0 \left(1 + \frac{A^2}{16} \right)$ et l'expression de $\epsilon \simeq -\frac{A^2}{192}$.

3. Étude d'un oscillateur à l'aide de son portrait de phase

On fait l'étude d'un oscillateur M de masse $m = 0,2 \text{ kg}$ astreint à se déplacer suivant l'axe Ox de vecteur unitaire $\vec{u}_x$. Il est soumis uniquement aux forces suivantes :

- la force de rappel d'un ressort de caractéristiques (k, ℓ_0) .
- une force de frottement visqueux : $\vec{f}_v = -\lambda \dot{x} \vec{u}_x$.
- une force constante $\vec{F}_c = F_c \vec{u}_x$

1. a) Établir l'équation différentielle du mouvement de M et la mettre sous la forme :

$$\ddot{x} + \frac{\omega_0}{Q} \dot{x} + \omega_0^2 x = \omega_0^2 X_0$$

où x est l'allongement du ressort (par rapport à ℓ_0). Les grandeurs ω_0 , Q et X_0 sont à exprimer en fonction des données.

b) Dans le cas d'une solution pseudo-périodique, exprimer $x(t)$: on définira le temps caractéristique τ et la pseudo-pulsation Ω que l'on exprimera en fonction de ω_0 et Q .

2. Le portrait de phase $(v(t) = \dot{x}(t), x(t))$ de l'oscillateur étudié est donné sur la figure (7.12). On souhaite pouvoir en tirer les valeurs des différents paramètres de l'oscillateur.

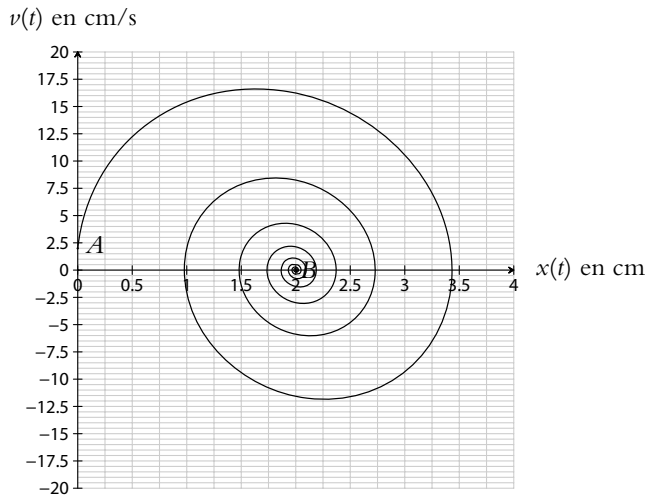


Figure 7.12

- Quel est le type de mouvement ?
- Déterminer la vitesse et l'élongation au début et à la fin du mouvement.
- Déterminer la vitesse maximale atteinte ainsi que l'élongation maximale.
- On donne les différentes dates correspondant aux croisements de la trajectoire de phase avec l'axe des abscisses :

$t \text{ (s)}$	0,31	0,65	0,97	1,3	1,62
-----------------	------	------	------	-----	------

En déduire la pseudo-période T et la pseudo-pulsation Ω .

e) On définit le décrétement logarithmique par $\delta = \frac{1}{n} \ln \left(\frac{x(t) - x_B}{x(t+nT) - x_B} \right)$ où $x(t)$ et $x(t+nT)$ sont les élongations aux instants t et $t+nT$ (n entier naturel) et x_B l'élongation finale de M . Exprimer δ en fonction de T et τ .

En choisissant une valeur de n la plus grande possible pour les données dont on dispose, déterminer δ puis τ .

f) Dédurre des résultats précédents le facteur de qualité Q et la pulsation propre ω_0 .

g) Déterminer la raideur du ressort k , le coefficient de frottement λ et la force F_c sachant que $\ell_0 = 1$ cm.

4. Modélisation d'un amortisseur

On considère l'amortisseur d'un véhicule. Chaque roue supporte un quart de la masse de la voiture assimilé à un point M de masse $m = 500$ kg, et est reliée à un amortisseur dont le ressort a une constante de raideur $k = 2,5 \cdot 10^4$ N.m⁻¹. Le point M subit aussi un frottement visqueux $\vec{f}_v = -\lambda \vec{v}$ où $\vec{v}$ est la vitesse de M et $\lambda = 5 \cdot 10^3$ kg.s⁻¹.

Le véhicule franchit à vitesse constante un défaut de la chaussée de hauteur $h = 5$ cm. Son inertie est suffisante pour qu'il ne se sou- lève pas immédiatement mais acquiert une

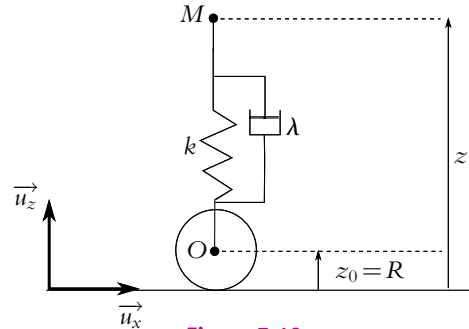


Figure 7.13

vitesse verticale $v_0 = 0,5$ m.s⁻¹. On pose $\alpha = \frac{\lambda}{2m}$. On note Z_i la cote du point M avant le passage du défaut.

1. a) On note $Z(t)$ la cote de M . Établir l'équation différentielle pour Z après le passage de l'obstacle. Déterminer $Z(t)$ en fonction des données. On remarquera que $\alpha = \Omega$ où Ω est la pseudo-pulsation.

b) Les passagers sont sensibles à l'accélération verticale de la voiture, calculer sa valeur maximale. On utilisera le fait que $\alpha = \Omega$.

2. Il faut en fait éviter des oscillations susceptibles de provoquer chez les passagers le mal des transports, en se plaçant dans les conditions critiques. Pour quelle masse par roue est-ce réalisé, k et λ étant inchangés ?

5. Oscillations avec frottement solide

On considère un mobile M de masse m se déplaçant selon un axe horizontal (Ox) et assimilé à un point matériel. Ce mobile est soumis à une force de rappel $-kx$ et à une force de frottement solide, le coefficient de frottement est f : la réaction $\vec{R}$ du support sur le point se décompose entre une composante tangentielle $\vec{R}_T$ et une composante normale $\vec{R}_N$, avec $\|\vec{R}_T\| = f \|\vec{R}_N\|$ et $\vec{R}_T \cdot \vec{v}(M)_{\text{support}} < 0$ si le mobile est en mouvement, et $\|\vec{R}_T\| < f \|\vec{R}_N\|$ si il est immobile.

On pose le mobile au point d'abscisse x_0 sans vitesse initiale.

1. Montrer que pour que le solide se mette en mouvement il faut que $|x_0| > a = \frac{fmg}{k}$.
2. On suppose $|x_0| > a$. Montrer que le mouvement du solide se décompose en plusieurs phases. Déterminer $x_1(t)$, abscisse du mobile au cours de la première phase du mouvement. À quel instant cette première phase s'arrête-t-elle ? À quelle condition le mobile s'arrête-t-il ?
3. On suppose que le mobile ne s'arrête pas. Donner la pseudo-période des oscillations du mobile. Comparer avec l'oscillateur amorti par frottement fluide étudié dans le cours.

6. Un modèle d'élasticité d'une fibre de verre, d'après ENSTIM 2004

Le verre est un matériau très dur. On peut toutefois le déformer légèrement sans le casser : on parle d'élasticité. Récemment, des expériences de biophysique ont été menées pour étudier l'ADN. Le capteur utilisé était simplement une fibre optique en silice amincie à l'extrémité de laquelle on accroche un brin d'ADN. L'expérience consistait à suivre la déformation de flexion de la fibre. La masse volumique du verre est $\rho = 2500 \text{ kg.m}^{-3}$.

La fibre de verre de longueur ℓ et de diamètre d est encastrée horizontalement dans une paroi immobile. Au repos, la fibre est horizontale (on néglige son poids). Quand on applique une force verticale F (on supposera que la force F reste verticale tout au long de l'expérience) à l'extrémité libre de la fibre, celle-ci est déformée. L'extrémité est déplacée verticalement d'une distance Y que l'on appelle la flèche (figure 7.14).

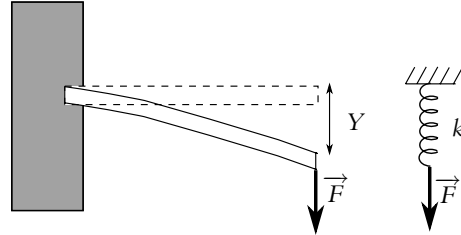


Figure 7.14

La flèche Y est donnée par la relation suivante (on notera la présence du facteur numérique 7, sans dimension, qui est en fait une valeur approchée pour plus de simplicité) : $Y = \frac{7\ell^3 F}{Ed^4}$, où E est appelé module d'Young du verre. Pour les applications numériques on prendra pour le module d'Young $E = 7.10^{10} \text{ S.I.}$

1. Quelle est l'unité S.I. du module d'Young E ?
2. En considérant uniquement la force F , montrer que l'on peut modéliser la fibre de verre par un ressort de longueur à vide nulle et de constante de raideur k dont on donnera l'expression analytique en fonction de E , d et ℓ .
3. Calculer numériquement k pour une fibre de longueur $\ell = 7 \text{ mm}$ et de diamètre $d = 10 \text{ }\mu\text{m}$.
4. Démontrer l'expression de l'énergie potentielle élastique d'un ressort de longueur à vide nulle, de constante de raideur k , lorsque sa longueur est ℓ . En reprenant l'analogie du ressort, quelle est alors l'énergie potentielle élastique de la fibre de verre lorsque la flèche vaut Y ?

On a tous fait l'expérience suivante : faire vibrer une règle ou une tige lorsque une de ses extrémités est bloquée. On cherche ici à trouver les grandeurs pertinentes qui fixent la fréquence des vibrations. L'extrémité de la tige vaut $Y(t)$ à l'instant t . On admet que lors des vibrations de la fibre, l'énergie cinétique de la fibre de verre est donnée par l'expression

$$E_c = \rho \ell d^2 \left(\frac{dY}{dt} \right)^2.$$

- Écrire l'expression de l'énergie mécanique de la fibre en négligeant l'énergie potentielle de pesanteur.
- Justifier que l'énergie mécanique se conserve au cours du temps. En déduire l'équation différentielle qui régit les vibrations de la fibre.
- Quelle est l'expression de la fréquence propre de vibration d'une tige de verre de module d'Young E , de longueur ℓ et de diamètre d .
- Calculer numériquement la fréquence des vibrations d'une fibre de verre de longueur 7 mm et de diamètre 0,01 mm.

7. Système de recul d'un canon

On considère un canon (figure 7.15) de masse $M = 800$ kg. Lors du tir horizontal d'un obus de masse $m = 2$ kg avec une vitesse $\vec{v}_0 = v_0 \vec{u}_x$ ($v_0 = 600$ m.s⁻¹), le canon acquiert une vitesse de recul $\vec{v}_c = -\frac{m}{M} \vec{v}_0$.

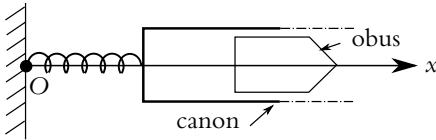


Figure 7.15

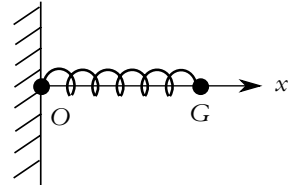


Figure 7.16

Pour limiter la course du canon, on utilise un ressort de raideur k_1 , de longueur à vide L_0 dont l'une des extrémités est fixe et l'autre lié au canon. Le déplacement a lieu suivant l'axe Ox . Dans la suite, le canon est assimilé à son centre de gravité G (figure 7.16).

- Quelle est la longueur du ressort lorsque le canon est au repos ?
- En utilisant l'énergie mécanique, déterminer la distance de recul d . En déduire la raideur k_1 pour avoir un recul au maximum égal à d . Application numérique pour $d = 1$ m.
- Retrouver la relation entre k_1 et d en appliquant la relation fondamentale de la dynamique.
- Quel est l'inconvénient d'utiliser un ressort seul ?

Pour pallier ce problème, on ajoute au système un dispositif amortisseur, exerçant une force de frottement visqueux $\vec{F} = -\lambda \vec{v}$, $\vec{v}$ étant la vitesse du canon.

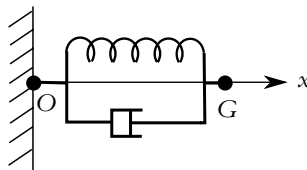


Figure 7.17

- Le dispositif de freinage absorbe une fraction $E_a = 778$ J de l'énergie cinétique initiale. Calculer la nouvelle valeur k_2 de la constante de raideur du ressort avec les données numériques précédentes. Déterminer la pulsation propre ω_0 de l'oscillateur.

6. Déterminer λ pour que le régime soit critique. Application numérique.
7. Déterminer l'expression de l'élongation $x(t)$ du ressort, ainsi que celle de la vitesse $\dot{x}(t)$. En déduire l'instant t_m pour lequel le recul est maximal. Exprimer alors ce recul en fonction de m , v_0 et λ . L'application numérique redonne-t-elle la valeur de d précédente ?



Optique

Approximation de l'optique géométrique, rayon lumineux

8

1. Notion expérimentale de rayon lumineux

On constate que la zone lumineuse issue d'une source ponctuelle, dans un milieu homogène, est située à l'intérieur d'un cône, appelé *faisceau lumineux*. Si on intercale sur le trajet de la lumière un écran percé d'un trou de petite taille, on restreint l'ouverture de la zone éclairée : on parle alors de *pinceau lumineux* (Cf. figure 8.1). L'image du pinceau sur un écran placé parallèlement au précédent écran sera un disque lumineux appelé *image géométrique* du pinceau.

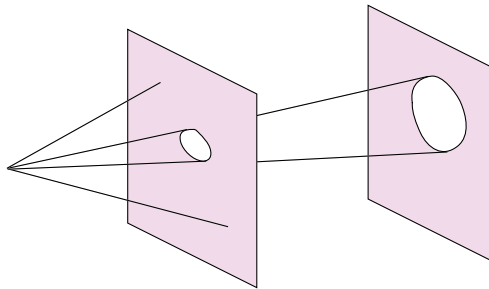


Figure 8.1 Faisceau et rayon lumineux.

Cependant, si on cherche à isoler une zone très étroite en diminuant fortement la taille du trou, pour observer précisément le trajet suivi par la lumière, c'est l'inverse qui se produit : **le pinceau s'élargit** et on observe de la lumière autour de l'image géométrique du pinceau. **Il est donc impossible d'isoler un rayon lumineux.** C'est une notion abstraite, servant pour les constructions ; en particulier, on utilisera graphiquement des rayons pour délimiter un faisceau.

On peut cependant imaginer à quoi pourrait ressembler un rayon lumineux en observant le faisceau très fin d'un laser.

Le phénomène, décrit ci-dessus, d'élargissement de la zone éclairée autour de l'image géométrique est caractéristique d'un comportement ondulatoire. Tout se passe comme si tous les points du trou se comportaient comme des sources secondaires, réémettant de la lumière dans toutes les directions. Cet éparpillement de la lumière est appelé *diffraction* : ce phénomène, présenté en terminale, sera approfondi en deuxième année. On peut s'en rendre compte en regardant une source lumineuse de petite taille (comme un réverbère au loin) à travers un rideau en voile fin. Au lieu d'une seule tache lumineuse, on observe d'autres taches lumineuses autour de la tache centrale.

Le phénomène de diffraction ne se rencontre pas qu'avec la lumière mais avec toute onde comme les ondes sonores par exemple.

Il se présente chaque fois que l'on limite le faisceau par une ouverture de très petites dimensions.

2. Aspect ondulatoire

Dans ce paragraphe, on rappelle et on complète certaines notions sur les ondes vues en terminale.

2.1 Phénomène ondulatoire

En 1665, R. Hooke émet l'idée que la lumière est une vibration de haute fréquence qui se propage. Cette hypothèse est développée par Huygens puis par Young au début du XIX^e siècle pour interpréter les phénomènes d'interférences et par Fresnel pour expliquer la diffraction. Entre-temps, Newton a émis l'idée d'une nature corpusculaire, qui sera reprise au XX^e siècle par Einstein et Planck. Maxwell, après avoir construit la théorie de l'électromagnétisme (1876), affirme que la lumière est une onde électromagnétique qui se propage dans le vide à la vitesse de 3.10^8 m.s^{-1} . Il précise que l'onde est transverse, c'est-à-dire que le champ électrique $\vec{E}$ et le champ magnétique $\vec{B}$ sont perpendiculaires à la direction de propagation (Cf. figure 8.2). C'est en fait au carré de la norme de $\vec{E}$ qu'est directement liée l'intensité lumineuse.

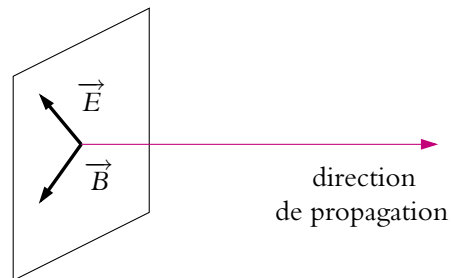


Figure 8.2 Propagation du champ électromagnétique.

On s'intéresse ici au cas particulier des ondes sinusoïdales polarisées rectilignement (c'est-à-dire telles que le vecteur champ électrique garde une direction constante), à partir desquelles on peut obtenir, par superposition, toute onde électromagnétique. Dans le cas d'une onde sinusoïdale qui se propage selon l'axe Ox , dans le sens des x croissants, l'expression du champ électrique en fonction du temps et de la position est

de la forme :

$$\vec{E} = \vec{E}_0 \cos \left(\omega \left(t - \frac{x}{v} \right) \right)$$

où v est la vitesse de propagation de l'onde. Ainsi à un instant donné, pour tous les points d'un plan perpendiculaire à l'axe Ox , la valeur du champ électrique est la même. On dit qu'ils appartiennent au même *plan d'onde*. De manière plus générale, l'ensemble des points où $\vec{E}$ prend la même valeur à un instant donné est appelé *surface d'onde*. Une onde lumineuse est dite *plane* si les surfaces d'onde sont des plans (ce qui est le cas donné précédemment) et *sphérique* si les surfaces d'onde sont des sphères.

L'énergie lumineuse se propage selon des courbes perpendiculaires en tout point aux surfaces d'onde (Cf. figures 8.3 et 8.4). Ce sont justement ces trajectoires de l'énergie lumineuse que l'on appelle *rayons lumineux*. Il s'agit donc d'une définition un peu différente du rayon lumineux de celle vue dans le paragraphe précédent.

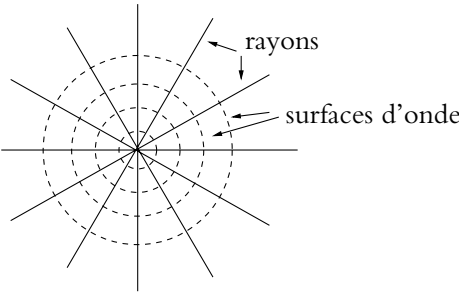


Figure 8.3 Onde sphérique.

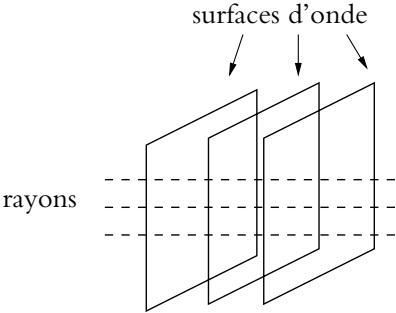


Figure 8.4 Onde plane.

2.2 Indice de réfraction d'un milieu

La vitesse de propagation (ou célérité) de l'onde dépend du milieu traversé. Dans un milieu autre que le vide, elle est inférieure à c :

$$v = \frac{c}{n} \quad \text{avec} \quad n > 1$$

La grandeur n est appelé *indice de réfraction du milieu*.

Voici quelques valeurs d'indice de réfraction :

Milieu	Indice
Vide	$n = 1$
Air (C.N.T.P.) ¹	$n = 1,00027 \simeq 1$
Eau	$n = 1,33$
Verre courant	$n \simeq 1,5$
Verre à fort indice	$1,6 \leq n \leq 1,8$

¹ Les Conditions Normales de Température et de Pression (C.N.T.P.) correspondent à une pression de 1 bar et une température de 273 K.

Pour les milieux transparents (c'est-à-dire non absorbants) et dans le visible, les indices sont supérieurs à 1. Dans la pratique, on trouve peu d'indice supérieur à 2 (l'indice du diamant, particulièrement élevé, est de l'ordre de 2,4).

2.3 Longueur d'onde

Tout phénomène périodique peut être représenté par une superposition d'ondes sinusoïdales, comme on le verra dans le cours d'électrocinétique de la deuxième période. Ainsi toute onde lumineuse peut être considérée comme une superposition d'ondes sinusoïdales comme celles décrites précédemment. D'après ce qui précède, la vitesse de propagation de l'onde dépend du milieu mais pas la pulsation ω .

Une onde comportant une seule pulsation est dite *monochromatique*, la pulsation étant caractéristique de la couleur du rayonnement lumineux.

On utilise aussi pour caractériser l'onde la période $T = \frac{2\pi}{\omega}$ ou la fréquence $\nu = \frac{1}{T}$. Une des grandeurs les plus utilisées est la *longueur d'onde* :

$$\lambda = \nu T = \frac{cT}{n} \quad (8.1)$$

grandeur qui dépend du milieu traversé. On peut exprimer la longueur d'onde λ dans un milieu en fonction de la longueur d'onde λ_0 dans le vide :

$$\lambda = \frac{\lambda_0}{n} \quad (8.2)$$

On remarquera que $\lambda_0 = cT$.

Les longueurs d'onde correspondant au domaine du visible varient de 400 nm à 750 nm pour un œil moyen. Les couleurs principales sont récapitulées dans le tableau suivant :

λ_0 (nm)	< 400	500	590	630	> 750
ν (Hz)	> 7,5 10 ¹⁴	6,0 10 ¹⁴	5,1 10 ¹⁴	> 4,8 10 ¹⁴	< 4,0 10 ¹⁴
Couleur	ultraviolet	bleu	jaune	rouge	infrarouge

On peut citer le cas du laser hélium-néon, de couleur rouge souvent utilisé en travaux pratiques dont la longueur d'onde est 632,8 nm.

Les indices de réfraction cités au paragraphe précédent dépendent de la longueur d'onde.

3. Lois fondamentales de l'optique géométrique

On considère un milieu *homogène* et *isotrope* dans toute la suite du cours. Mais que signifient les termes homogène et isotrope ?

Homogène signifie que les propriétés physiques (composition, densité, indice de réfraction...) sont les mêmes en tout point du milieu.

Isotrope signifie que ces propriétés sont identiques dans toutes les directions de propagation du rayon lumineux.

Par exemple, une autoroute est un milieu homogène pour les voitures car elles peuvent s'y déplacer à la même vitesse quel que soit l'endroit ; par contre, elles ne peuvent se déplacer que dans une direction et deux sens donc ce n'est pas un milieu isotrope. Il existe des milieux en optique pour lesquels la vitesse de propagation n'est pas la même parallèlement aux trois axes de coordonnées : on dit que ces milieux sont *biréfringents* (c'est par exemple le cas du quartz). Ils ne sont pas étudiés dans cet ouvrage.

3.1 Cadre de l'approximation de l'optique géométrique

Dans certaines conditions, le caractère ondulatoire de la lumière est peu apparent, c'est-à-dire que des phénomènes caractéristiques comme la diffraction ne sont pas observables. C'est le cas quand toutes les ouvertures traversées sont grandes par rapport à la longueur d'onde. On retrouve les observations décrites au début du chapitre. On dit alors qu'on se trouve dans le cadre de l'approximation de l'optique géométrique, ce qui entraîne alors un certain nombre de propriétés simples énoncées dans la suite.

Les lois de l'optique géométrique étudiées dans la suite sont valables tant que les instruments utilisés sont de grande taille par rapport à la longueur d'onde.

3.2 Propagation rectiligne

Depuis la plus haute Antiquité, il est observé que la lumière issue d'une source ponctuelle se propageant dans un milieu homogène et isotrope emprunte un trajet rectiligne passant par la source (on peut remarquer que, à indice constant donc à vitesse constante, cela revient à minimiser le temps de parcours).

Dans un milieu homogène et isotrope, les rayons lumineux sont des droites.

Ainsi la prévision du comportement de la lumière pourra se faire directement à l'aide des lois de la géométrie, et c'est pour cela que l'on parle d'optique géométrique.

Dans une suite de milieux homogènes, le trajet d'un rayon lumineux sera formé d'une succession de segments de droite.

3.3 Retour inverse

Soient deux points A et B situés sur un rayon lumineux. On constate que si on inverse le sens de propagation de la lumière (en changeant la source lumineuse de place par exemple), A et B sont toujours sur le même rayon lumineux.

La trajectoire suivie par la lumière ne dépend pas du sens de parcours.

Ceci constitue la loi du retour inverse de la lumière.

3.4 Indépendance des rayons lumineux

Soient deux sources de lumière S_1 et S_2 . On observe que les rayons lumineux provenant de S_1 et S_2 ne se « gênent pas », c'est-à-dire que la répartition lumineuse obtenue est la somme des répartitions individuelles. On dit qu'il y a **indépendance des rayons lumineux**.

Réfraction, réflexion

9

1. Le vocabulaire de l'optique géométrique

1.1 Dioptries et miroirs

On appelle *système optique* l'ensemble d'un certain nombre de milieux transparents (c'est-à-dire non absorbants) en général homogènes et isotropes séparés par des surfaces dont la forme est simple (portions de plans ou de sphères le plus souvent). Si la surface entre deux milieux successifs est réfléchissante, on parle de *miroir* ; sinon il s'agit d'un *dioptre*.

Lorsqu'un système ne contient que des dioptries, on parle de *système dioptrique* ; s'il contient un ou plusieurs miroirs, on parle de *système catadioptrique*.

1.2 Axe optique, plans transverses

Le plus souvent, le système est centré c'est-à-dire qu'il possède un axe de symétrie, dit *axe optique*. On oriente l'axe optique dans le sens de propagation de la lumière. Quand il y a réflexion, l'orientation change de sens. Les plans perpendiculaires à cet axe optique sont appelés *plans transverses*.

1.3 Image

Soit un système optique (S) et soit une source ponctuelle de lumière placée en A . Si toute la lumière issue de A converge en A' (Cf. figure 9.1) après avoir traversé (S), on dit que A' est l'*image* de A à travers (S) et que A est l'*antécédent* de A' . Les points A et A' sont dits *conjugués* par le système. Par le retour inverse de la lumière, si l'on place la source ponctuelle en A' , éclairant le système (S), toute la lumière issue de A' et passant par (S) converge en A .

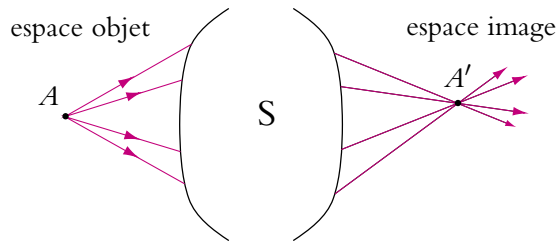


Figure 9.1 Espaces objet et image.

1.4 Faces d'entrée et de sortie

Le système est limité par deux surfaces. La face d'entrée est la face sur laquelle la lumière arrive, la face de sortie est celle d'où la lumière repart (elles sont confondues pour un miroir). Le système (S) divise l'espace en deux : l'*espace objet* situé en avant de la face d'entrée, l'*espace image* situé en arrière de la face de sortie. Pour un miroir, l'espace objet et l'espace image sont géométriquement confondus.



Pour définir l'espace objet et l'espace image, il faut toujours se référer au sens de propagation de la lumière.

1.5 Réalité et virtualité

a) Faisceau lumineux

Un rayon *incident* se dirige vers la face d'entrée, alors qu'un rayon *émergent* quitte la face de sortie.

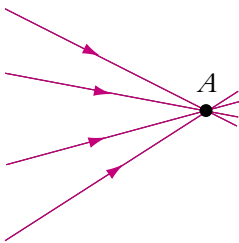


Figure 9.2 Faisceau convergent.

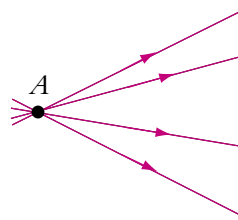


Figure 9.3 Faisceau divergent.

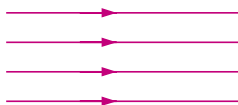


Figure 9.4 Faisceau parallèle.

Il existe trois types de faisceau qui sont représentés sur les figures 9.2, 9.3, 9.4. Un faisceau est :

- *convergent* si tous les rayons convergent vers un même point ;
- *divergent* si tous les rayons sont issus d'un même point ;
- *parallèle* si tous les rayons sont parallèles entre eux. Dans ce cas les rayons sont issus d'un point à l'infini et vont converger vers un point lui aussi à l'infini.

b) Objet réel ou virtuel, image réelle ou virtuelle

Si le faisceau incident diverge à partir d'un point A situé dans l'espace objet, on dit que A est un *objet réel*. Si, par contre, le faisceau diverge à partir d'un point A qui n'est pas situé dans l'espace objet, A est un *objet virtuel*.

De la même façon, si le faisceau émergent converge vers un point A' de l'espace image, A' est une *image réelle*. Sinon, A' est une *image virtuelle*. Les quatre cas possibles d'objet et d'image sont récapitulés sur les figures 9.5 à 9.8. On représente en pointillés les prolongement des rayons car ils ne correspondent pas à des trajets réellement suivis par la lumière.

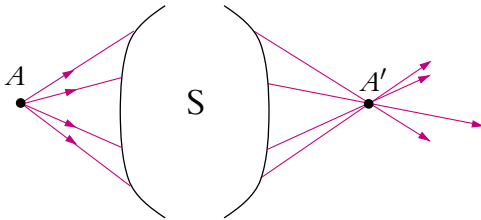


Figure 9.5 Objet réel - Image réelle.

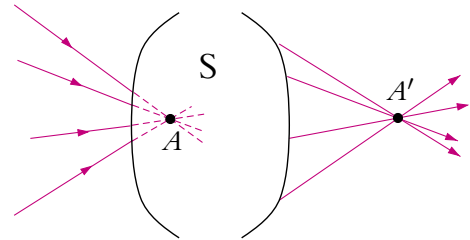


Figure 9.6 Objet virtuel - Image réelle.

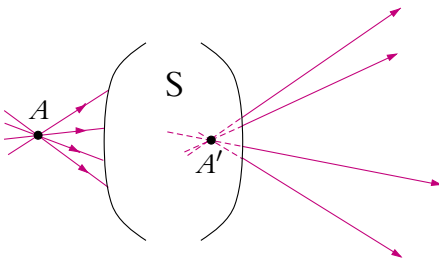


Figure 9.7 Objet réel - Image virtuelle.

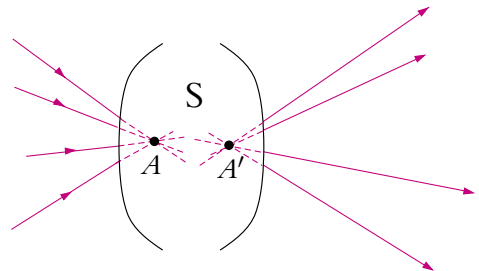


Figure 9.8 Objet virtuel - Image virtuelle.



Pratiquement, l'objet est réel si on peut le toucher, l'image est réelle si on peut la recueillir sur un écran sans utiliser un autre système optique.

On réalise facilement ce que représentent un objet réel et une image réelle car ce sont des notions concrètes. Qu'en est-il d'un objet virtuel ou d'une image virtuelle ? Tout le monde a déjà utilisé une loupe ou observé les yeux d'un interlocuteur qui porte des lunettes. Dans ces cas, l'observateur a l'impression de voir des objets derrière la loupe ou derrière les lunettes. La figure 9.9 représente le cas de la loupe. Les rayons issus de l'objet réel A sont déviés par la loupe. Après la loupe, le faisceau semble issu de A' qui est l'image de A par la loupe. **Cette image est virtuelle car ce ne sont pas les rayons mais leur prolongement qui passent réellement par A' .** Cependant pour l'œil qui observe, les rayons semblent provenir de A' et non de A : il verra un objet en A' .

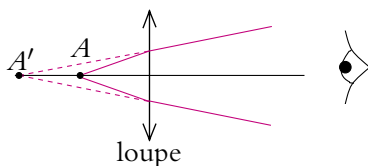


Figure 9.9 Loupe.

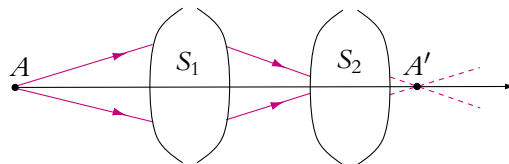


Figure 9.10 Formation d'un objet virtuel.

Comment fabriquer un objet virtuel ? Pour cela, il faut obtenir une image réelle A' d'un objet A avec un premier système optique (S_1), puis interposer un deuxième système optique (S_2) entre A' et (S_1). Alors A' est un objet virtuel pour (S_2) comme cela est représenté sur la figure 9.10.

Si le faisceau est parallèle, cela revient à rejeter l'objet ou l'image à l'infini.

2. Lois de Descartes

2.1 Réflexion

La réflexion consiste en un brusque changement de direction de la lumière incidente qui, après avoir rencontré une surface réfléchissante, revient dans son milieu de propagation initial.

Si la surface réfléchissante est en métal poli, on parle de *réflexion métallique*. Le facteur de réflexion en énergie, c'est-à-dire la fraction de la puissance de la lumière incidente qui est réfléchi, est alors proche de 1 (il dépend de l'inclinaison des rayons lumineux et aussi fortement de la longueur d'onde). Il faut noter que les miroirs domestiques comportent une plaque de verre protectrice sous laquelle est déposée la couche réfléchissante. Cela peut donner des images parasites comme on le verra en exercice. Une surface réfléchissante est représentée avec des hachures.

Si la surface réfléchissante est la surface de séparation entre deux milieux transparents, on parle de *réflexion vitreuse* (ou dioptrique). Le facteur de réflexion en énergie dépend des indices des deux milieux ainsi que de l'inclinaison du rayon lumineux. Bien qu'on

en tient rarement compte, un dioptré réfléchit toujours une partie de la lumière comme le lecteur a pu s'en rendre compte en observant des reflets dans une vitre.

Le point I où le rayon incident rencontre la surface est appelé *point d'incidence*. Le plan contenant le rayon incident et la normale à la surface au point I est appelé *plan d'incidence* (Cf. figure 9.11).

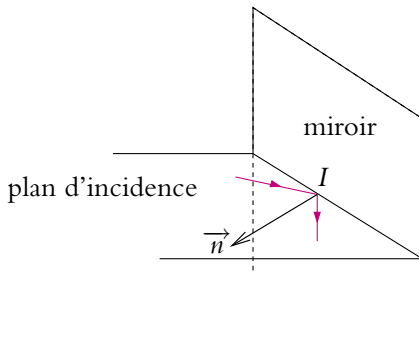


Figure 9.11 Réflexion sur un miroir.

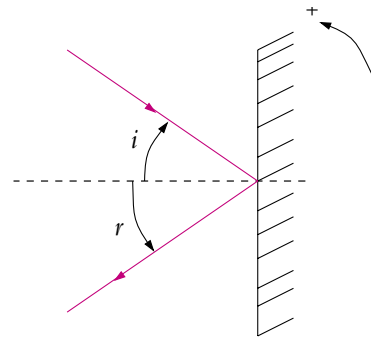


Figure 9.12 Représentation de la réflexion dans le plan d'incidence.

Les lois de la réflexion sont les suivantes :

1. le rayon réfléchi est dans le plan d'incidence,
2. les angles d'incidence i et de réflexion r sont tels que :

$$r = -i \quad (9.1)$$

Étant donné que les rayons incident et réfléchi sont dans le même plan, on a l'habitude de ne représenter que ce plan (Cf. figure 9.12). Il faut **choisir un sens d'orientation pour les angles** qui sont des grandeurs algébriques. Ici c'est le sens trigonométrique qui a été choisi, mais on peut choisir le sens horaire si cela est plus aisé. En revanche, **il faut absolument orienter les angles i et r à partir de la normale au miroir.**

2.2 Réfraction

La réfraction consiste en un brusque changement de direction de la lumière incidente qui, après avoir rencontré une surface dite *réfractante* ou *réfringente*, se propage dans un milieu différent de son milieu de propagation initial. On appellera n_1 l'indice de réfraction du milieu initial (dit aussi milieu incident) et n_2 l'indice du milieu final.

a) Lois de la réfraction

On définit le plan d'incidence comme au paragraphe précédent. Les lois de la réfraction ont été établies par W. Snell (physicien hollandais) en 1621 ; R. Descartes les a

retrouvées en 1637, avec celles de la réflexion, si bien qu'elles sont connues en France sous le nom de lois de Descartes. Elles sont les suivantes :

1. le rayon réfracté est dans le plan d'incidence,
2. les angles d'incidence i_1 et de réfraction i_2 sont tels que :

$$n_1 \sin i_1 = n_2 \sin i_2.$$

Étant donné que les rayons incident et réfracté sont dans le plan d'incidence, les schémas se font dans ce plan (Cf. figures 9.13 et 9.14). Comme précédemment, les angles sont orientés à partir de la normale.

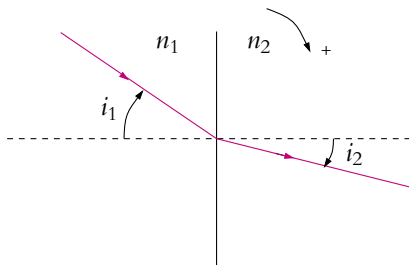


Figure 9.13 Réfraction avec un milieu (2) plus réfringent.

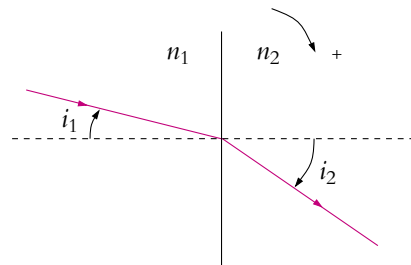


Figure 9.14 Réfraction avec un milieu (2) moins réfringent.

La relation entre les angles est donc :

$$n_1 \sin i_1 = n_2 \sin i_2 \quad (9.2)$$

Si $n_2 > n_1$, on dit que le milieu (2) est plus réfringent que le milieu (1). Dans ce cas, comme $\sin i_2 = \frac{n_1}{n_2} \sin i_1$, $\sin i_2 < \sin i_1$ et $i_2 < i_1$: **le rayon réfracté se rapproche de la normale** (Cf. figure 9.13). C'est le cas par exemple lorsque le milieu initial est l'air et le milieu final le verre ou l'eau. Cela explique le fait qu'un bâton plongé dans l'eau apparaisse cassé. On peut penser aussi au pêcheur utilisant un harpon. Il a l'impression de voir le poisson dans le prolongement des rayons lumineux arrivant dans son œil, alors que l'animal est plus proche (Cf. figure 9.15). S'il n'a aucune expérience, il aura tendance à lancer le harpon à un mauvais endroit. On peut se rassurer en pensant qu'en général le poisson est assez gros et que le pêcheur visant la tête touchera le corps !

Si $n_2 < n_1$, on dit que le milieu (2) est moins réfringent que le milieu (1). Dans ce cas, $\sin i_2 > \sin i_1$ et $i_2 > i_1$: **le rayon réfracté s'écarte de la normale** (Cf. figure 9.14). On retrouve d'ailleurs cette propriété à partir du cas précédent et du retour inverse de la lumière.

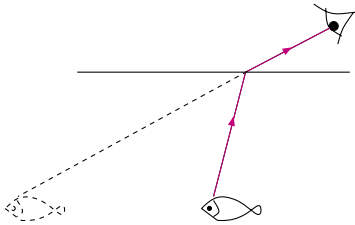


Figure 9.15 Influence de la réfraction sur la vision.

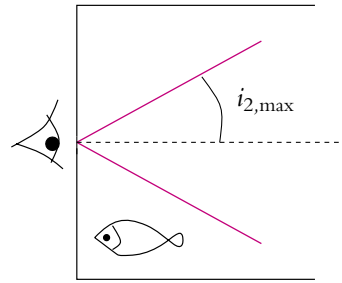


Figure 9.16 Cône de réfraction.



Lorsqu'on effectue un schéma, il faut donc prendre garde de dessiner correctement les rayons en fonction des indices.

b) Réflexion totale

Dans le cas où $n_2 < n_1$ (Cf. figure 9.14), si on augmente progressivement l'angle d'incidence i_1 , le rayon réfracté s'écarte de plus en plus de la normale jusqu'à valoir $\frac{\pi}{2}$. Dans ce cas si on continue à augmenter i_1 , il ne peut plus y avoir réfraction puisque le rayon réfracté ne serait plus dans le milieu (2). Il existe donc un angle limite R_{lim} pour i_1 , correspondant à $i_2 = \frac{\pi}{2}$:

$$\sin R_{\text{lim}} = \frac{n_2}{n_1} \quad (9.3)$$

Ainsi, pour $i_1 > R_{\text{lim}}$, comme cela vient d'être écrit, il ne peut plus y avoir de rayon réfracté, mais seulement un rayon réfléchi : il y a *réflexion totale*. Ce phénomène peut être observé par exemple par un nageur sous l'eau qui, en se retournant, regarde la surface et ne voit pas l'extérieur.

c) Réfraction limite

Dans le cas où $n_2 > n_1$, le rayon se rapproche de la normale donc le problème de la réflexion totale ne se pose pas. Les angles i_1 et i_2 vérifient $i_2 < i_1$: la valeur maximale de i_2 correspond à $i_1 = \frac{\pi}{2}$ et $\sin i_{2,\text{max}} = \frac{n_1}{n_2}$. Le rayon émergent est contenu dans un cône de sommet le point d'incidence, de demi-angle au sommet $i_{2,\text{max}}$, appelé *cône de réfraction*. Ainsi, si on imagine un aquarium sur lequel on a placé un cache percé d'un petit trou, on ne peut voir à travers ce trou que ce qui se trouve à l'intérieur du cône de réfraction. On ne voit pas, par exemple, le poisson de la figure 9.16. Dans ce cas, on n'a pas tenu compte de la traversée de la paroi en verre par les rayons.

► **Remarques :**

- Lorsque l'angle d'incidence est égal à $\frac{\pi}{2}$, on parle d'incidence rasante. Dans la réalité, l'angle est légèrement inférieur à $\frac{\pi}{2}$, sinon la lumière ne pourrait pas traverser le dioptre.
- Par retour inverse de la lumière, $i_{2,\max}$ devient R_{lim} .

2.3 Construction de Descartes

Il s'agit d'une méthode de construction précise des rayons réfractés représentée sur la figure 9.17. On cherche à exprimer géométriquement la relation de Descartes.

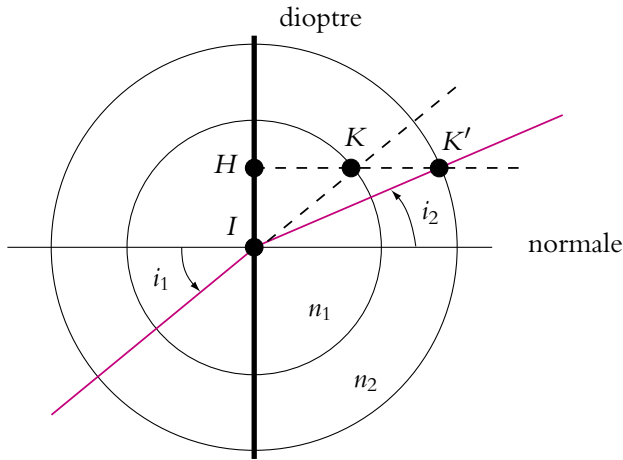


Figure 9.17 Construction de Descartes.

On trace deux cercles concentriques de centre I (point d'incidence), de rayon $R_1 = n_1$ et $R_2 = n_2$, n_1 et n_2 correspondant aux indices des deux milieux. En fait, les rayons R_1 et R_2 des cercles doivent simplement être dans le rapport des indices.

Sur la figure 9.17, le dioptre est représenté par le trait vertical passant par le point I et la normale au dioptre par le trait horizontal. On trace le rayon incident passant par I et faisant l'angle i_1 avec la normale. On prolonge ce rayon jusqu'à l'intersection avec le cercle de rayon n_1 , au point K . La projection de K sur le dioptre donne le point H tel que $IH = IK \sin i_1$ soit avec $IK = n_1$, $IH = n_1 \sin i_1$. Si on trace la parallèle à la normale passant par H , elle croise le cercle de rayon n_2 en K' . La projection de IK' sur le dioptre est aussi égale à IH . Donc l'angle i_2 que fait IK' avec la normale est tel que $IH = IK' \sin i_2 = n_2 \sin i_2$. Les points K et K' sont tels que $n_1 \sin i_1 = n_2 \sin i_2$ donc IK' représente le rayon réfracté.

3. Cas des milieux stratifiés. Mirages

3.1 Réfraction dans un milieu stratifié

Un milieu stratifié est un empilement de plans d'indice différent (Cf. figure 9.18). Sur la figure, on a choisi des indices croissant vers le bas, soit $n_1 > n_6$. Un rayon partant du bas et arrivant sur le premier dioptre avec l'angle d'incidence i_1 est réfracté avec l'angle $i_2 > i_1$ puisque $n_1 > n_2$.

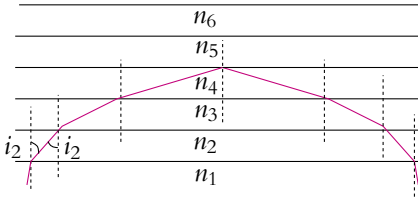


Figure 9.18 Milieu stratifié.

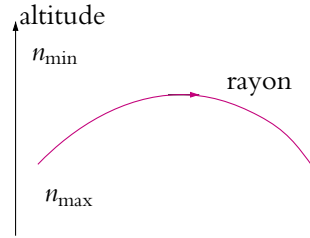


Figure 9.19 Milieu continu.

Ainsi, à chaque passage de dioptre le rayon s'écarte de la normale jusqu'à ce qu'il y ait réflexion totale sur l'un des dioptres. Le rayon repart alors vers le bas. **La trajectoire de la lumière est alors formée de segments de droites dont l'ensemble a sa concavité tournée vers les indices croissants** (ici vers le bas).

L'angle i_2 est aussi l'angle d'incidence sur le deuxième dioptre. En écrivant la relation de Descartes pour chaque dioptre, on obtient :

$$n_1 \sin i_1 = n_2 \sin i_2 = n_3 \sin i_3 = n_4 \sin i_4 = n_5 \sin i_5 \dots$$

On peut donc en déduire que dans un milieu stratifié composé de plans, les rayons vérifient :

$$n \sin i = \text{constante} \quad (9.4)$$

Dans un milieu continu dont l'indice varie, l'épaisseur des couches tend vers 0 et la courbe s'arrondit. Ainsi, dans un milieu dont l'indice croît vers le bas, on obtient le trajet lumineux de la figure 9.19. Ce phénomène a des conséquences entre autres en astronomie. En effet, les astronomes doivent tenir compte de la traversée de l'atmosphère par les rayons lumineux pour diriger leur télescope dans la bonne direction. Quand on observe une étoile (Cf. figure 9.20), on a l'impression qu'elle est dans la direction des rayons qui arrivent dans l'œil, on voit une position apparente alors que l'étoile est en fait plus basse sous l'horizon. L'atmosphère étant moins dense en haute altitude, l'indice de réfraction décroît avec l'altitude selon la loi de Gladstone $n - 1 = k\rho$, où n est l'indice de l'air, k une constante et ρ la masse volumique. On se trouve donc dans la situation de la figure 9.19. Cela aurait peu d'incidence sur les observations si l'étoile ne se déplaçait pas dans le ciel. Au fur et à mesure

qu'elle se déplace dans le ciel, la couche d'atmosphère traversée par la lumière varie et donc l'angle dont a tourné le rayon varie aussi. Les logiciels de pilotage des télescopes doivent tenir compte de ce phénomène.

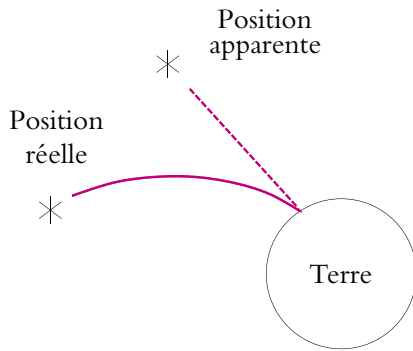


Figure 9.20 Position apparente d'une étoile.

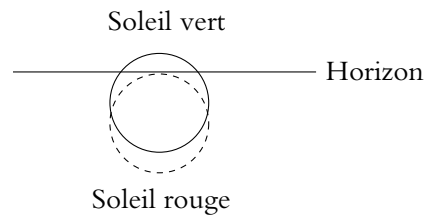


Figure 9.21 Rayon vert.

La réfraction atmosphérique a aussi une incidence sur la hauteur apparente du Soleil. Ainsi, on voit le soir le Soleil à l'horizon alors qu'il est déjà « couché » et passé sous l'horizon. L'indice de l'air comme celui d'autres milieux dépend de la longueur d'onde. Il suit la loi de Cauchy :

$$n = a + \frac{b}{\lambda^2}$$

où a et b sont des constantes positives. Pour l'eau, l'effet est moins important que pour le verre mais, dans tous les cas, le rouge est moins dévié que le vert. Ainsi lorsque le Soleil est sous l'horizon, il arrive qu'on observe une parcelle d'image verte du Soleil alors que l'image rouge moins déviée est complètement en dessous de l'horizon (Cf. figure 9.21). Ce phénomène est appelé *le rayon vert*.

3.2 Mirages

Tout le monde a déjà eu l'impression en été en circulant en voiture qu'il y avait de l'eau au loin sur la route. Le phénomène de mirage est un phénomène réel et non une divagation de l'esprit. Il est dû à la variation importante de température des couches d'air au voisinage du sol en fonction de l'altitude. Lorsque la température augmente, l'air se dilate et sa masse volumique ρ diminue (voir cours de thermodynamique). Or, d'après la loi de Gladstone citée dans le paragraphe précédent, l'indice de réfraction diminue si ρ diminue. Si on s'intéresse à une route dont le revêtement (noir) est surchauffé par le Soleil, l'air au voisinage de la route est chauffé par la route et la température décroît avec l'altitude. On obtient les variations T , n et ρ représentées sur la figure 9.22 en fonction de l'altitude z .

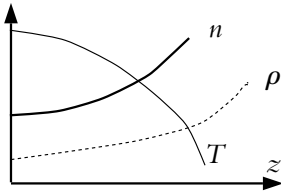


Figure 9.22 Variation de la température T , de l'indice n et de la masse volumique ρ de l'atmosphère terrestre en fonction de l'altitude.

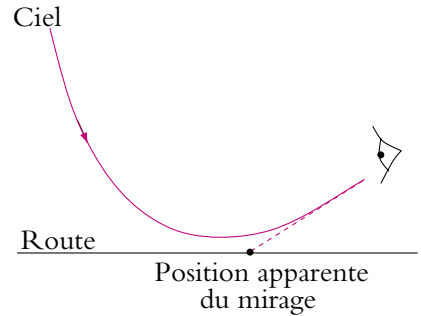


Figure 9.23 Mirage inférieur.

Puisque, lorsqu'on s'approche de la route l'indice de réfraction décroît, d'après l'étude du paragraphe précédent, un rayon provenant du ciel va se courber en s'approchant du sol avec sa concavité tournée vers le haut (Cf. figure 9.23). Lorsque le rayon pénètre dans l'œil de l'observateur, il semble provenir de la route (direction en pointillés sur la figure), et l'observateur aura l'impression de voir le ciel (ou plutôt une tache d'eau) sur la route au loin. C'est le même phénomène qui se produit dans le désert. Ce mirage porte le nom de *mirage inférieur* car on le voit sur le sol.

Il arrive que l'on assiste à des phénomènes de *mirages supérieurs* c'est-à-dire dans le ciel : c'est le contraire du mirage inférieur.

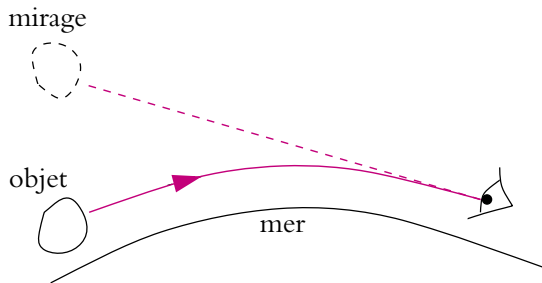


Figure 9.24 Mirage supérieur.

L'auteur a assisté à la formation d'un mirage supérieur sur la côte Vendéenne, en face de l'île de Ré. À une certaine heure de la journée, l'air au-dessus de la mer est plus chaud que l'eau et la température dans l'air croît avec l'altitude. Dans ce cas, les rayons sont courbés vers le sol (Cf. figure 9.24) et on peut voir une image de l'île de Ré renversée au dessus de l'île réelle.

A. Applications directes du cours

1. Erreur sur le positionnement d'un objet

Un observateur de taille $t = 1,8$ m regarde un objet au fond d'un bassin depuis le bord. Le bassin de hauteur $h = 1,5$ m contient une épaisseur $e = 1$ m d'eau (indice $n = 1,33$). L'objet semble situé à une distance $d = 1$ m du bord. Déterminer la véritable position de l'objet.

2. Panneau réfléchissant

Certains panneaux d'autoroute ont des inscriptions qui paraissent lumineuses dans l'éclairage des phares. Pour cela on peut utiliser des petites billes de fort indice de réfraction n (l'indice de l'air est pris égal à 1). On note R le rayon de la bille.

On considère un rayon incident parallèle à un axe de la bille, à une distance $h \leq R/2$ de cet axe (figure 9.25). Ce rayon subit une première réfraction, puis une réflexion sur la face opposée et enfin une seconde réfraction.

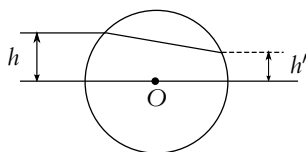


Figure 9.25

1. Quelle erreur commet-on sur la première réfraction en se plaçant dans l'approximation de Gauss (voir chapitre suivant, paragraphe 2.3) ? Dans la suite on se place dans cette approximation.
2. À quelle distance h' de l'axe le rayon frappe-t-il la deuxième face de la bille ? En déduire la condition sur n pour avoir $h' = 0$
3. En déduire la direction d'émergence du rayon après la deuxième réfraction.

3. Détection de pluie sur un pare-brise

On modélise un pare-brise par une lame de verre à faces parallèles, d'épaisseur e , d'indice $n_v = 1,5$. Un rayon lumineux issu d'un détecteur E (figure 9.26) arrive sur le dioptre en I avec un angle d'incidence i .

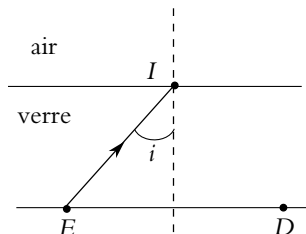


Figure 9.26

1. a) Quelle est la condition sur i pour qu'il y ait réflexion totale en I ? On suppose dans la suite $i = 60^\circ$.
- b) Où faut-il placer le détecteur de lumière D (figure 9.26)?
2. Dans le cas de pluie, une lame d'eau (épaisseur e') se dépose sur le pare-brise. L'indice de l'eau est $n_e = 1,33$.
- a) Déterminer le trajet du rayon lumineux.
- b) En déduire le fonctionnement d'un détecteur de pluie.

4. Prisme isocèle rectangle argenté

Soit un prisme rectangle isocèle dont on a argenté l'hypothénuse. Des rayons arrivent perpendiculairement à cette face. On note x et y les distances du sommet du prisme à l'axe des rayons respectivement incident et émergent correspondant.

Montrer que la distance $x + y$ est indépendante du rayon incident considéré.

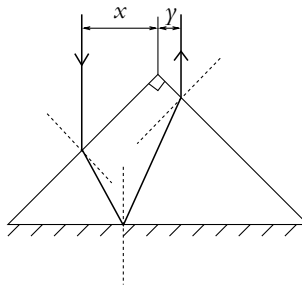


Figure 9.27

5. Incidence de Brewster

Un rayon lumineux arrive à l'interface plane séparant l'air d'un milieu d'indice n . Il se scinde en un rayon réfléchi et un rayon réfracté. On souhaite obtenir que ces deux rayons soient perpendiculaires. Déterminer la valeur des différents angles.

Application numérique : $n = 1,33$.

B. Exercices et problèmes

1. Prisme de renvoi d'un vidéoprojecteur, d'après ESSAIM 2005

En vue d'obtenir une grande compacité des vidéoprojecteurs, on utilise un prisme d'entrée (PE) et un prisme de sortie (PS) accolés (il y a en fait une mince couche d'air entre les deux prismes) pour renvoyer la lumière de la lampe d'éclairage vers une surface composée de micromiroirs. Cette surface doit renvoyer la lumière selon l'axe optique de la lentille de projection.

Le plan des miroirs est incliné d'un angle $\alpha = 10^\circ$ par rapport aux faces JK et JJ' .

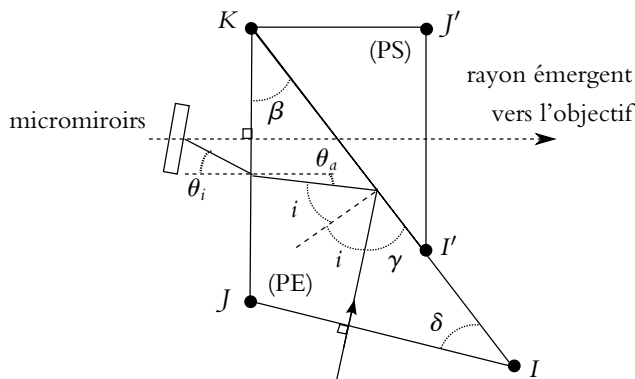


Figure 9.28

1. Calculer l'angle limite pour lequel il y a réflexion totale dans le cas d'une interface verre ($n = 1,5$)-air ($n = 1$).
2. a) Déterminer la valeur de l'angle d'arrivée θ_a du faisceau sur la face (JK) pour que le rayon soit réfléchi par les miroirs selon l'axe optique.
b) Exprimer θ_a en fonction des angles i et β seuls.
c) Quelle est la condition sur i pour que le faisceau incident se réfléchisse bien sur la face (IK)? Quelle valeur limite β_1 de β en déduit-on?
d) Quelle est l'autre valeur limite β_2 de β permettant que le faisceau réfléchi sur les miroirs traverse bien l'interface (IK)?
e) On choisit d'imposer $i = 45^\circ$, calculer β et δ .

2. Réfractomètre de Pulfrich et d'Abbe

Un réfractomètre est un dispositif permettant de mesurer les indices de matériau solide ou liquide.

On considère un prisme d'angle au sommet A et d'indice n . L'une des faces est placée horizontalement et on pose sur cette face un matériau d'indice N . On suppose que le contact est parfait au sens où l'interface est plane et où il n'y a pas d'air entre le prisme et le matériau.

1. Lorsqu'un rayon lumineux passe d'un milieu d'indice n_1 à un milieu d'indice $n_2 < n_1$, à quelle condition sur l'angle d'incidence i_1 obtient-on un rayon réfracté?
2. Donner les conditions nécessaires pour observer le phénomène de réflexion totale.
3. Représenter le dispositif et tracer le chemin suivi par un rayon lumineux arrivant à l'interface entre le matériau et le prisme avec un angle d'incidence i . On note r l'angle de réfraction s'il existe, β l'angle d'incidence à la surface entre le prisme et l'air et α l'angle de réfraction correspondant s'il existe.
4. La direction de la surface d'entrée (entre l'air et le matériau) du dispositif a-t-elle une importance? On se demandera si tous les rayons peuvent entrer dans le dispositif.

5. Donner en la justifiant la condition sur les indices permettant l'absence de réflexion totale à l'interface entre le matériau et le prisme quelles que soient les circonstances.
6. Préciser le domaine de variation de i puis de r .
7. Exprimer β en fonction de A et r et en déduire le domaine de variation de β .
8. Discuter avec précision la possibilité de réflexion totale entre le prisme et l'air. La discussion portera sur la valeur de A .
9. En déduire les conditions pour avoir un rayon émergeant du dispositif.
10. Pour le réfractomètre de Pulfrich, on a $A = 90^\circ$ et $n = 1,732$. Donner, dans ce cas, les valeurs de β permettant l'émergence des rayons.
11. Qu'observe-t-on lorsqu'on regarde sous différents angles la face de sortie du réfractomètre (face entre le prisme et l'air)? On introduira une valeur particulière de l'angle α sous lequel on regarde la face.
12. La valeur limite de α prenant la valeur de 30° , en déduire l'expression de N en fonction de n et de α et donner sa valeur numérique.
13. Pour le réfractomètre d'Abbe, on a $A = 61^\circ$ et $n = 1,60$. Donner, dans ce cas, les valeurs de β permettant l'émergence des rayons.
14. Donner la relation entre N , n et la valeur limite de α introduite pour le réfractomètre de Pulfrich.
15. La valeur limite de α prenant la valeur de 15° , en déduire la valeur numérique de N . On précisera le détail de la résolution numérique.
16. Que se passe-t-il si on introduit une goutte d'un liquide d'indice $n' > N$?
17. Même question pour une goutte d'un liquide d'indice $n' < N$?
18. Conclure sur les méthodes à adopter pour mesurer l'indice d'un solide ou d'un liquide.

3. Télémètre

Un télémètre est un appareil permettant de mesurer une distance à l'aide d'une approche optique.

On considère un dispositif constitué de deux miroirs, l'un totalement réfléchissant et l'autre semi-réfléchissant. Le premier M_2 est un miroir classique tandis que le second M_1 permet au rayon d'être réfracté lorsqu'il arrive sur la surface non réfléchissante. Les deux miroirs sont inclinés d'un angle φ l'un par rapport à l'autre. On place une source ponctuelle S du côté de la face non réfléchissante du miroir M_1 à une distance L du centre O_1 de M_1 . Un rayon incident arrive sur M_1 en O_1 avec une incidence de 45° et traverse M_1 . Un autre rayon arrive sur M_2 avec une incidence α au centre O_2 de M_2 , il est réfléchi sur M_2 puis sur la surface réfléchissante de M_1 . On tourne M_2 jusqu'à superposer les deux images. On note d la distance de O_2 aux rayons émergeant du système.

1. En assimilant M_1 à une lame à faces parallèles lorsqu'il est traversé par les rayons, montrer qu'alors le rayon émergent est parallèle au rayon incident.

2. On note β l'angle entre la direction du rayon SO_2 et celle du miroir semi-réfléchissant M_1 . Etablir la relation entre β et φ .
3. En déduire la relation entre φ et γ l'angle entre SO_1 et SO_2 .
4. Montrer que ce dispositif permet de déterminer L si on connaît d et la valeur de φ pour que les deux images se superposent.

4. Dispersion de la lumière solaire lors d'un arc-en-ciel, d'après CCP 2005

Dans l'ensemble de ce problème, tous les angles seront considérés en valeur arithmétique, c'est-à-dire réels positifs. Leurs valeurs numériques seront à exprimer en degrés décimaux.

1. Préciser exactement ce que l'on entend par dispersion.

2. Questions préliminaires.

- a) Rappeler les lois de Descartes pour la réfraction d'un rayon lumineux passant de l'air (milieu d'indice unité) vers un milieu d'indice n . On fera un schéma en notant i l'angle d'incidence et r l'angle de réfraction. Exprimer la dérivée $\frac{dr}{di}$ exclusivement en fonction de l'indice n et du sinus ($\sin i$) de l'angle d'incidence.
- b) Exprimer, en fonction de i et de r , la valeur de la déviation du rayon lumineux, définie par l'angle entre la direction incidente et la direction émergente, orientées dans le sens de propagation.
- c) Exprimer aussi, à l'appui d'un schéma, la déviation d'un rayon lumineux dans le cas d'une réflexion.

3. Lorsque le soleil illumine un rideau de pluie, on peut admettre que chaque goutte d'eau se comporte comme une sphère réceptionnant un faisceau de rayons parallèles entre eux.

Dans tout ce qui suit, on considérera que l'observation est faite par un oeil accommodant à l'infini, c'est-à-dire assimilable à une lentille convergente (cristallin) capable de focaliser sur un écran (rétine) tout faisceau de lumière parallèle issu d'une goutte d'eau.

On recherche, dans un premier temps, les conditions pour que la lumière émergente, issue d'une goutte d'eau, se présente sous forme d'un faisceau de lumière parallèle. Pour cela on fait intervenir l'angle de déviation D de la lumière à travers la goutte d'eau, mesuré entre le rayon émergent et le rayon incident. Cet angle de déviation D est une fonction de l'angle d'incidence i .

Exprimer la condition de parallélisme des rayons émergents en la traduisant mathématiquement au moyen de la dérivée $\frac{dD}{di}$.

4. Une goutte d'eau quelconque, représentée par une sphère de centre O et de rayon R , est atteinte par la lumière solaire sous des incidences variables, comprises entre 0° et 90° . Son indice, pour une radiation donnée, sera noté n tandis que celui de l'air sera pris égal à l'unité.

Dans chacun des trois cas suivants :

- 1) Lumière directement transmise (Figure 9.29),

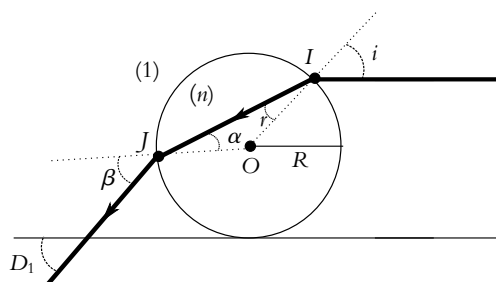


Figure 9.29

2) Lumière transmise après une réflexion partielle à l'intérieur de la goutte (Figure 9.30),

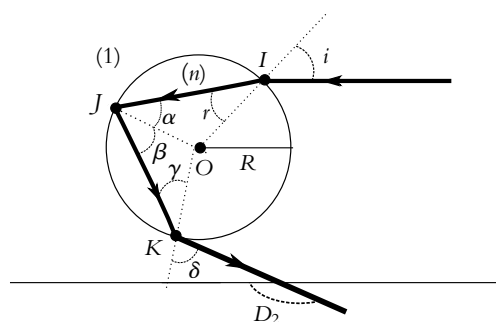


Figure 9.30

3) Lumière transmise après deux réflexions partielles à l'intérieur de la goutte (Figure 9.31),

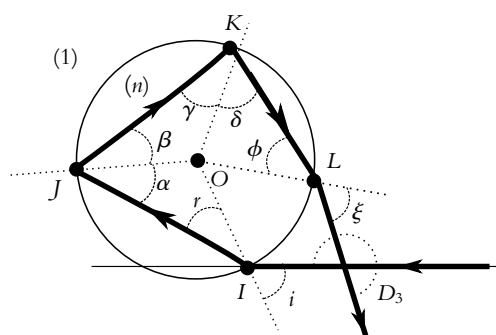


Figure 9.31

répondre successivement aux questions a, b, c, ci-après :

a) Exprimer en fonction de l'angle d'incidence i ou de l'angle de réfraction r , tous les angles marqués de lettres grecques.

- b)** En déduire l'angle de déviation D propre à chaque cas, en fonction de i et de r .
- c)** Rechercher ensuite, si elle existe, une condition d'émergence d'un faisceau parallèle, exprimée par une relation entre le sinus ($\sin i$) de l'angle d'incidence et l'indice n de l'eau.
- 5.** Le soleil étant supposé très bas sur l'horizon, normal au dos d'un observateur, montrer que celui-ci ne pourra observer la lumière transmise que si la goutte d'eau se trouve sur deux cônes d'axes confondus avec la direction solaire et de demi-angles au sommet $\theta_2 = 180^\circ - D_2$ (justification de l'arc primaire) et $\theta_3 = D_3 - 180^\circ$ (justification de l'arc secondaire).
- 6.** Les angles θ_2 et θ_3 dépendant de l'indice n de l'eau, on observe un phénomène d'irisation dû au fait que cet indice évolue en fonction de la longueur d'onde.

Calculer ces angles pour le rouge et le violet, sachant que pour le rouge l'indice vaut 1,3317 tandis que pour le violet il est égal à 1,3448.

En admettant que l'observateur se trouve face à un rideau de pluie, dessiner la figure qui apparaît dans son plan d'observation en notant la position respective des rouges et des violets.

Stigmatisme et aplanétisme. Dioptries et miroirs

10

1. Stigmatisme et aplanétisme rigoureux

1.1 Définitions

a) Stigmatisme rigoureux

Un système optique (S) est dit *rigoureusement stigmatique* pour le couple de points (A , A') si tous les rayons issus de A passent par A' après avoir été déviés par le système. Les points A et A' sont dits *conjugués* par rapport à (S).

L'image d'un point de l'axe par un système centré est forcément sur l'axe. En effet, comme l'axe optique est un axe de symétrie, la normale au système au point d'incidence sur l'axe est colinéaire à l'axe. Ainsi un rayon colinéaire à l'axe n'est pas dévié car il est confondu avec la normale.

b) Aplanétisme rigoureux

Soient deux points A et A' de l'axe optique conjugués par rapport à (S). Soit B , un point du plan transverse passant par A . Le système (S) sera dit *aplanétique* pour A et A' si le conjugué de B , noté B' , se trouve dans le plan transverse passant par A' . Il y a alors correspondance plan transverse par plan transverse.

1.2 Notion de foyer

Soit un objet très éloigné de (S), situé sur l'axe : il est dit « à l'infini ». Son image à travers (S) est appelée *foyer image* F' et se trouve sur l'axe optique. Le plan transverse contenant F' est appelé *plan focal image* (Cf. figure 10.1). Réciproquement, le *foyer objet* F est le point de l'axe optique dont l'image est rejetée à l'infini sur l'axe (Cf. figure 10.2). Le plan transverse contenant F est appelé *plan focal objet*.

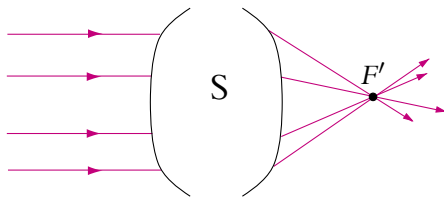


Figure 10.1 Foyer image.

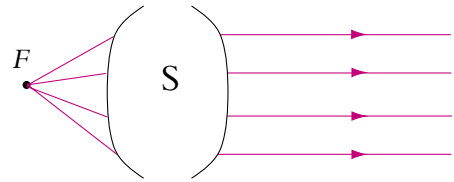


Figure 10.2 Foyer objet.

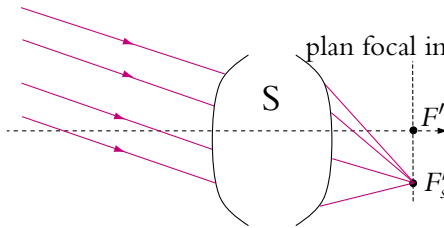


Figure 10.3 Foyer image secondaire.

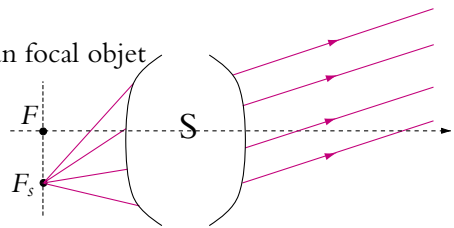


Figure 10.4 Foyer objet secondaire.

Soit un objet très éloigné, hors de l'axe optique. On a déjà signalé que les rayons incidents étaient parallèles entre eux mais pas forcément parallèles à l'axe optique. Pour un système aplanétique, l'image sera dans le même plan que l'image d'un objet à l'infini sur l'axe, donc dans le plan focal image. Le point image F'_s est alors appelé *foyer image secondaire* (Cf. figure 10.3).

De même, un point F_s du plan focal objet, dit *foyer objet secondaire*, donnera une image à l'infini hors de l'axe (Cf. figure 10.4).

1.3 Exemple de système rigoureusement stigmatique : le miroir plan

Il faut déjà se poser la question suivante : comment, connaissant la position de l'objet, peut-on déterminer la position de l'image ? On a vu que l'image est à l'intersection de tous les rayons (ou de leur prolongement) provenant de l'objet et passant par le système. Il suffit donc de construire deux de ces rayons et l'image sera à leur intersection si elle est réelle ou à l'intersection de leur prolongement si elle est virtuelle.

Soit A un objet réel. On réalise la construction (Cf. figure S10.8) de deux rayons en utilisant les lois de Descartes sur le miroir. Les rayons réfléchis ne se croisent pas mais **ils semblent provenir de A'** . Un observateur qui regarde dans le miroir aura l'impression de voir l'image en A' . Le point A' est à l'intersection des prolongements des rayons réfléchis. **A' est l'image virtuelle, symétrique de A par rapport au miroir.**

Tout point de l'espace objet aura ainsi un conjugué, symétrique par rapport au miroir plan, qui est donc stigmatique pour tout point de l'espace.

En utilisant le retour inverse de la lumière et en considérant A comme un objet virtuel, on obtient avec la même construction A' image réelle de A (Cf. figure S10.9).

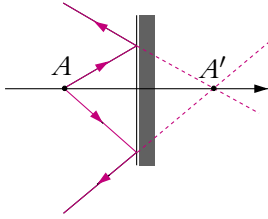


Figure 10.5 Image d'un objet réel par un miroir plan.

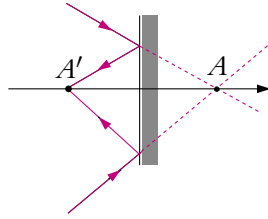


Figure 10.6 Image d'un objet virtuel par un miroir plan.

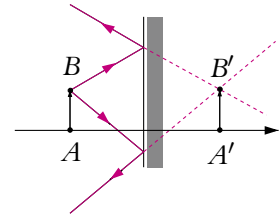


Figure 10.7 Aplanétisme d'un miroir plan.

Soit B dans le plan transverse contenant A . Par la même construction translatée vers le haut (Cf. figure 10.7), on obtient B' symétrique de B par rapport au miroir, et donc dans le plan transverse contenant A' .

Le miroir plan est rigoureusement aplanétique pour tout point de l'espace et c'est le seul système optique qui vérifie cette propriété.

On peut donc construire l'image d'un objet linéique transverse, AB , à travers le miroir ; tout point du segment AB sera sur le segment $A'B'$. L'image a la même taille que l'objet. On définit le **grandissement**, qui est une grandeur algébrique, par :

$$\gamma = \frac{A'B'}{AB} \quad (10.1)$$

Ici $\gamma = 1$, ce qui signifie que l'image est de même taille que l'objet et dans le même sens. $A'B'$ est symétrique de AB par le miroir. Comme objet et image sont symétriques par rapport au miroir, l'image d'un point à l'infini est à l'infini : **les foyers objet et image d'un miroir plan sont à l'infini**. On dit que le système est *afocal*.

2. Stigmatisme et aplanétisme approchés

2.1 Stigmatisme approché

Comme on va le voir sur de nombreux exemples, l'image d'un point n'est pas un point mais une tache lumineuse. Peut-on se contenter de la tache obtenue ? Oui, si elle n'est pas trop grosse. Mais que signifie « pas trop grosse » ? En fait, cela dépend de l'observateur, qui utilise soit son œil, soit un autre récepteur tel qu'une pellicule photo ou un récepteur numérique (plaque CCD). Dans le cas de l'œil, ce sont les cellules (cônes ou bâtonnets) qui assurent la réception (Cf. chapitre sur les

instruments de travaux pratiques), une pellicule est formée de grains plus ou moins gros tandis qu'une barrette CCD est formée elle aussi de petites cellules de réception.

Dans tous les cas, un récepteur est composé d'une multitude de petits récepteurs élémentaires qui ont une surface S .

Le récepteur ne peut distinguer des détails plus petits que S . En effet, si deux photons arrivent sur le même récepteur élémentaire, l'observateur verra une seule tache lumineuse. Ainsi, si deux photons provenant du même objet arrivent sur la même cellule élémentaire du récepteur après avoir traversé l'instrument optique, l'observateur aura l'impression de voir seulement un point ; par contre, si les photons arrivent sur deux cellules différentes, l'observateur verra une tache.

En conclusion, un instrument optique peut donner d'un objet ponctuel une image non ponctuelle, à condition que la tache image obtenue soit de taille plus petite que la cellule élémentaire du détecteur. Dans ce cas, l'observateur aura l'impression de voir un point. On parlera alors de *stigmatisme approché*.

De même, on pourra parler d'aplanétisme approché car les images d'objets dans le même plan seront vus quasiment dans le même plan par l'observateur.

2.2 Cas du dioptré plan

On s'intéresse à l'exemple simple d'un dioptré plan séparant deux milieux d'indice n et n' , avec $n' < n$, pour voir quelles sont les conditions du stigmatisme approché. Dans ce cas de choix d'indice, la réfraction éloigne les rayons de la normale. On considère un point A de l'axe et on recherche son image par le dioptré. Comme cela a été fait avec le miroir, il s'agit de trouver l'intersection de deux rayons après réfraction. Étant donné que le point A est sur l'axe, le point image A' est aussi sur l'axe (un rayon colinéaire à l'axe n'est pas dévié par le dioptré). Par contre, suivant l'inclinaison du rayon incident, l'intersection du rayon émergent avec l'axe est différente car plus le rayon incident arrive sur le dioptré avec un angle important, plus l'angle après réfraction sera grand (Cf. figure 10.8).

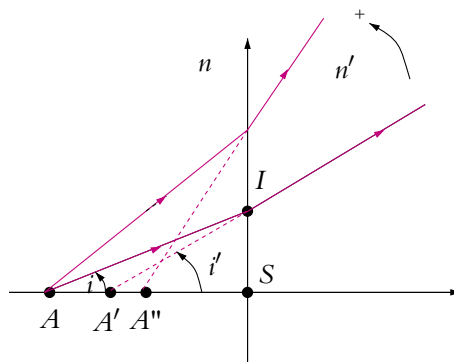


Figure 10.8 Stigmatisme d'un dioptré plan.

Il n'y a donc pas stigmatisme rigoureux.

On cherche à quelles conditions on pourra considérer qu'il y a stigmatisme approché c'est-à-dire à quelles conditions A' est quasiment indépendant du rayon.

Il s'agit donc de calculer la position de l'image A' . On utilise les relations dans les triangles AIS et $A'IS$, soit :

$$\overline{SI} = \overline{AS} \tan i \quad \text{et} \quad \overline{SI} = \overline{A'S} \tan i'$$

d'où :

$$\overline{A'S} = \overline{AS} \frac{\sin i \cos i'}{\sin i' \cos i}$$

On utilise alors la loi de Descartes $n \sin i = n' \sin i'$ ainsi que les relations $\cos i = \sqrt{1 - \sin^2 i}$ et $\cos i' = \sqrt{1 - \sin^2 i'}$, ce qui donne :

$$\overline{A'S} = \overline{AS} \frac{n'}{n} \sqrt{\frac{1 - \left(\frac{n}{n'}\right)^2 \sin^2 i}{1 - \sin^2 i}}$$

Dans le cas où $\sin i \simeq 0$, la relation précédente devient :

$$\frac{\overline{SA'}}{n'} = \frac{\overline{SA}}{n} \quad (10.2)$$

Le résultat est indépendant du rayon si i est très petit, c'est-à-dire pour des rayons très peu inclinés sur l'axe : on dit, dans ce cas, que le dioptré est utilisé dans les conditions de Gauss.

Les conditions de Gauss, qui viennent d'être évoquées, seront développées dans le paragraphe suivant.

La relation (10.2) est appelée *relation de conjugaison* du dioptré plan utilisé dans les conditions de Gauss avec origine au sommet S .

Comme pour le miroir, on peut effectuer une translation dans un plan transverse pour trouver l'image d'un point B voisin de A dans ce plan transverse : B' sera dans le plan transverse contenant A' et $\overline{A'B'} = \overline{AB}$. Il y a donc aplanétisme approché : dans les conditions de Gauss, l'image d'un plan est un plan.

Le grandissement transversal est égal à :

$$\gamma = \frac{\overline{A'B'}}{\overline{AB}} = 1$$

Si l'objet A est à l'infini soit $\overline{SA} \rightarrow \infty$, $\overline{SA'} \rightarrow \infty$ d'après la relation (10.2). L'image d'un point à l'infini est à l'infini donc, comme le miroir plan, **le dioptré plan est afocal.**

2.3 Conditions de Gauss, optique paraxiale

Les conditions du stigmatisme et de l'aplanétisme approchés sont apparues dans l'étude du dioptré plan. On peut aussi évoquer l'exemple des écrans de cinéma, qui doivent être courbes et non plus plans lorsqu'on souhaite obtenir des grandes images de

manière à ce que tous les points de l'écran soient à même distance du projecteur. Cela signifie que l'aplanétisme approché n'est plus vérifié.

En plus de ces constatations, on peut s'intéresser aux résultats de simulations de tracé de rayons lumineux dans le cas d'un miroir sphérique qui, comme on le verra dans le paragraphe suivant, est une portion de sphère réfléchissante. La première figure (figure 10.9) représente des rayons issus d'une source ponctuelle et réfléchis sur l'ensemble du miroir. Il apparaît clairement qu'il n'existe pas de point d'intersection unique des rayons réfléchis : **l'image d'un point n'est pas un point** et il n'y pas stigmatisme rigoureux. On ne peut même pas parler de stigmatisme approché car les rayons se croisent en des endroits très différents.

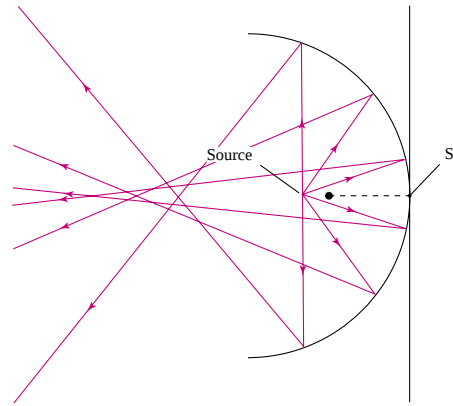


Figure 10.9 Miroir sphérique en dehors des conditions de Gauss.

En revanche, si l'on sélectionne les rayons se réfléchissant près du sommet S du miroir, on remarque qu'ils se croisent après réflexion en un point : on peut alors considérer qu'il y a stigmatisme. La figure 10.10 illustre cette propriété. Sur cette figure, on a agrandi la zone du miroir autour du sommet, c'est pour cela que le miroir n'apparaît que comme un arc de cercle. **On retiendra donc que les rayons doivent frapper le miroir au voisinage de son sommet.**

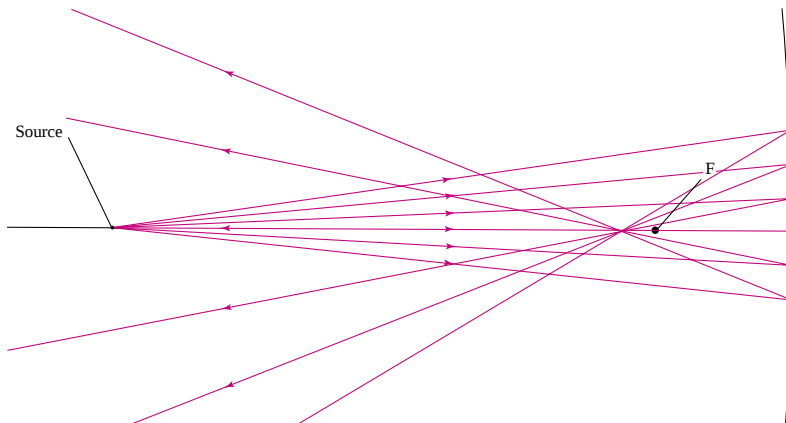


Figure 10.10 Miroir sphérique dans les conditions de Gauss.

La propriété précédente n'est pas suffisante. En effet, comme le montre la figure 10.11, si les rayons sont très inclinés par rapport à l'axe, il n'y a plus stigmatisme.

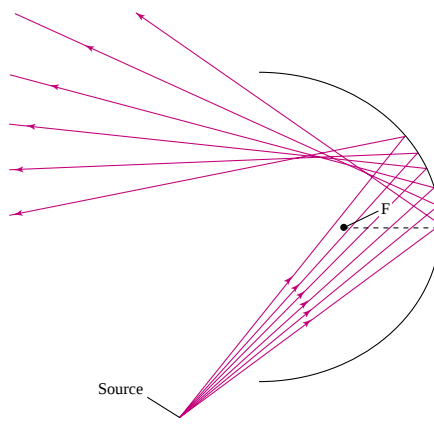


Figure 10.11 Miroir sphérique en dehors des conditions de Gauss.

Ces constatations amènent à énoncer certaines conditions pour que le stigmatisme et l'aplanétisme approchés soient vérifiés : on les appelle les *conditions de Gauss*.

Elles sont au nombre de trois mais deux seulement sont nécessaires, la troisième est une conséquence des deux autres.

1. Les rayons lumineux font un angle petit avec l'axe du système. On parle de *rayons paraxiaux*.
2. Les rayons rencontrent les dioptries ou les miroirs au voisinage de leur sommet situé sur l'axe optique si le système est centré.
3. L'angle d'incidence des rayons sur les dioptries ou les miroirs est petit.

Cela signifie aussi que l'objet et l'image ne doivent pas être de taille trop grande.

On va tirer les conséquences de ces conditions sur différents systèmes étudiés dans la suite.

3. Cas du miroir sphérique

3.1 Relation de Thalès dans les triangles

On rappelle ici les relations de Thalès dans les triangles, souvent utilisées en optique géométrique.

Dans les triangles représentés (Cf. figure 10.12), on peut écrire les relations de Thalès :

$$\frac{\overline{DE}}{\overline{AB}} = \frac{\overline{CD}}{\overline{CA}} = \frac{\overline{CE}}{\overline{CB}} \quad (10.3)$$

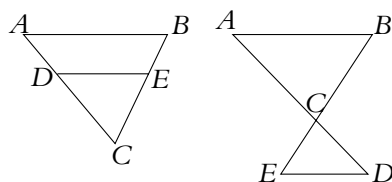


Figure 10.12 Relation de Thalès.

3.2 Miroirs sphériques

Un *miroir sphérique* est une portion de sphère réfléchissante. Il existe deux types de miroirs : le miroir *concave* (Cf. figure 10.13) et le miroir *convexe* (Cf. figure 10.14).

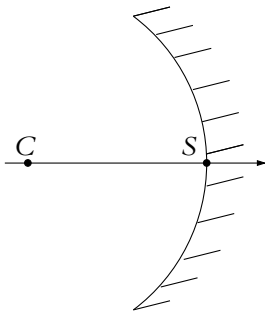


Figure 10.13 Miroir concave.

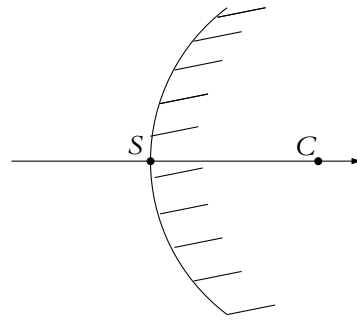


Figure 10.14 Miroir convexe.

On définit :

- le *centre* du miroir C qui est le centre de la sphère ;
- le *sommet* du miroir S qui est l'intersection du miroir avec l'axe optique ;
- le *rayon algébrique* du miroir $R = \overline{CS}$. Il est positif pour le miroir concave, négatif pour le miroir convexe.

Le miroir plan est un cas particulier du miroir sphérique pour lequel le rayon est infini.

3.3 Stigmatisme et aplanétisme approchés du miroir sphérique

Les simulations présentées dans le paragraphe sur les conditions de Gauss ont montré le stigmatisme approché pour le miroir. Qu'en est-il de l'aplanétisme ? On recherche grâce au logiciel de tracé de rayons l'image d'un point B dans le plan transverse d'un point A (Cf. figure 10.15). On s'aperçoit que si l'objet est petit (conditions de Gauss), l'image B' est dans le plan transverse de l'image A' : **il y a donc aplanétisme approché**.

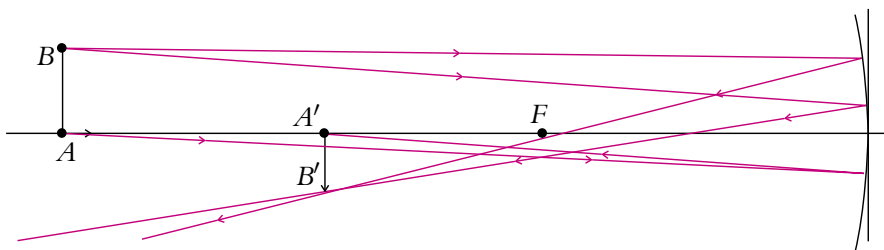


Figure 10.15 Aplanétisme approché d'un miroir sphérique.

3.4 Caractère focal des miroirs sphériques dans l'approximation de Gauss. Distance focale

Si on envoie un faisceau de rayons parallèles à l'axe sur un miroir concave, dans les conditions de Gauss, les rayons réfléchis se croisent sur l'axe au point focal image F' qui, d'après la simulation de la figure 10.16, est à mi-distance du centre C et du sommet S du miroir. Par retour inverse de la lumière, on conclut que tout rayon issu de F' se réfléchira parallèlement à l'axe et que F' est aussi le foyer objet du miroir.

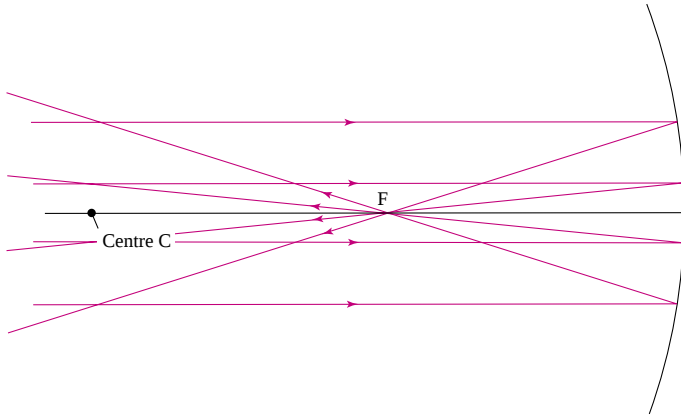


Figure 10.16 Foyer d'un miroir sphérique.

La figure 10.16 montre le cas d'un miroir concave. Dans ce cas, le point $F = F'$ est réel. Pour un miroir convexe, on a aussi $F = F'$ mais les foyers sont virtuels.

Le miroir sphérique comporte deux foyers confondus. Ils sont réels pour un miroir concave, virtuels pour un miroir convexe (Cf. figure 10.17) .

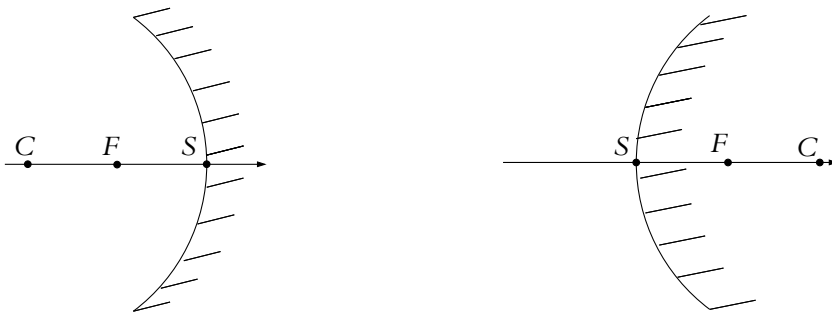


Figure 10.17 Position des foyers d'un miroir sphérique.

On définit la *distance focale objet* f par $f = \overline{SF}$ et la *distance focale image* f' par $f' = \overline{SF'}$. Ces deux grandeurs sont égales et valent :

$$f = f' = \frac{\overline{SC}}{2}$$

Puisqu'il y a aplanétisme approché, on peut définir les plans focaux objet et image qui sont confondus : il s'agit du plan transverse passant par $F = F'$.

3.5 Représentation symbolique des miroirs

Pour montrer que c'est dans les conditions de Gauss qu'on utilise les miroirs, on ne les représente pas par un arc de cercle mais par les symboles de la figure 10.18. En fait, comme cela apparaît sur la simulation de la figure 10.10, puisqu'on ne s'intéresse qu'aux rayons se réfléchissant au voisinage du centre S du miroir, on représente le miroir par son plan tangent en S .

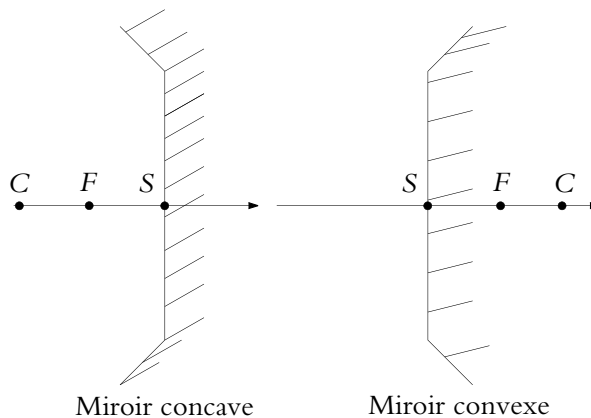


Figure 10.18 Symbole des miroirs sphériques.

3.6 Quelques règles de construction préliminaires

À partir des propriétés énoncées dans les paragraphes précédents, on va établir les relations de conjugaison des miroirs permettant de déterminer les positions des objets et des images ainsi que les grandissements.

Ces relations sont établies dans le cas particulier d'un miroir concave avec un objet placé avant le centre C du miroir mais sont valables quelle que soit la position de l'objet ainsi que pour un miroir convexe car elles sont algébriques.

On se sert de trois règles de construction qui seront revues plus en détail dans la suite du chapitre :

- un rayon passant par le centre du miroir n'est pas dévié puisqu'il frappe le miroir selon la normale ;
- un rayon incident parallèle à l'axe donne un rayon réfléchi passant par le foyer image ;
- un rayon incident passant par le foyer objet donne un rayon réfléchi parallèle à l'axe.

Ces trois règles permettent d'effectuer la construction de l'image $A'B'$ de l'objet AB sur la figure 10.19.

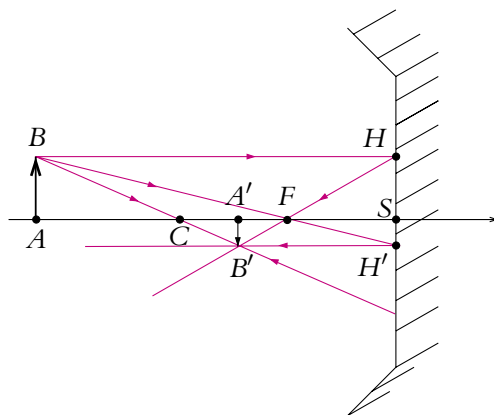


Figure 10.19 Construction d'une image par un miroir concave.

► **Remarque :** Pour obtenir des formules faisant intervenir le centre optique, il faut tracer un rayon passant par C . De même, pour retrouver des formules faisant intervenir le foyer, il faut tracer les rayons passant par F .

3.7 Relations avec origine aux foyers, dites relations de Newton

Sur la figure 10.19, on applique les relations de Thalès après avoir construit les points H et H' projections de B et B' sur le miroir. Dans les triangles BAF et FSH' :

$$\gamma = \frac{\overline{A'B'}}{\overline{AB}} = \frac{\overline{SH'}}{\overline{AB}} = \frac{\overline{FS}}{\overline{FA}}$$

Dans les triangles $B'A'F$ et FSH :

$$\gamma = \frac{\overline{A'B'}}{\overline{AB}} = \frac{\overline{A'B'}}{\overline{SH}} = \frac{\overline{FA'}}{\overline{FS}}$$

En égalant les deux expressions précédentes, on obtient la relation de conjugaison avec origine aux foyers, appelée relation de Newton :

$$\overline{FA} \overline{FA'} = \overline{FS}^2 = f^2 = ff' \quad (10.4)$$

3.8 Relation avec origine au centre

Comme précédemment, l'étude est faite pour le miroir concave mais toutes les expressions sont valables pour le miroir convexe. Les notations font référence à la figure 10.19.

Pour obtenir le grandissement $\gamma = \frac{\overline{A'B'}}{\overline{AB}}$, il suffit d'appliquer la relation de Thalès dans les triangles CAB et $CA'B'$. On obtient :

$$\gamma = \frac{\overline{A'B'}}{\overline{AB}} = \frac{\overline{CA'}}{\overline{CA}} \quad (10.5)$$

Pour établir la relation de conjugaison avec origine en C , on repart de la relation de Newton (10.4) et on introduit le point désiré C . Cela donne :

$$(\overline{FC} + \overline{CA})(\overline{FC} + \overline{CA'}) = f^2$$

Or $\overline{FC} = f$ d'où en développant :

$$f^2 + f \overline{CA} + f \overline{CA'} + \overline{CA} \overline{CA'} = f^2$$

On simplifie et on divise l'expression par $f \overline{CA} \overline{CA'}$:

$$\frac{1}{\overline{CA'}} + \frac{1}{\overline{CA}} = -\frac{1}{f}$$

Sachant que $f = \frac{\overline{SC}}{2}$, on obtient alors la formule de conjugaison du miroir sphérique dite avec origine au centre C :

$$\frac{1}{\overline{CA}} + \frac{1}{\overline{CA'}} = \frac{2}{\overline{CS}} \quad (10.6)$$

► **Remarque** : D'après la relation précédente, on remarque que les points S et C sont leur propre image par le miroir.

3.9 Formules de conjugaison avec origine au sommet

Il est parfois plus facile d'utiliser une relation de conjugaison avec origine au sommet S . On procède comme pour la formule de conjugaison avec origine en C : on introduit le point S par la relation de Chasles dans l'expression (10.4).

$$(\overline{FS} + \overline{SA})(\overline{FS} + \overline{SA'}) = f^2$$

Or $\overline{FS} = -f$ d'où en développant :

$$f^2 - f \overline{SA} - f \overline{SA'} + \overline{SA} \overline{SA'} = f^2$$

On simplifie et on divise l'expression par $f \overline{SA} \overline{SA'}$:

$$\frac{1}{\overline{SA'}} + \frac{1}{\overline{SA}} = \frac{1}{f}$$

Sachant que $f = \frac{\overline{SC}}{2}$, on obtient la formule de conjugaison avec origine au sommet :

$$\frac{1}{\overline{SA}} + \frac{1}{\overline{SA'}} = \frac{2}{\overline{SC}} \quad (10.7)$$



Les deux formules de conjugaison se ressemblent ; cependant il faut se souvenir que ou bien C est toujours en premier dans les grandeurs algébriques, ou bien c'est S.

Pour obtenir la formule de grandissement avec origine au sommet, il est nécessaire de tracer un rayon passant par S (Cf. figure 10.20).

Pour déterminer la position de B' , on utilise le rayon issu de B passant par C et celui issu de B passant par S. Le premier n'est pas dévié comme cela à déjà été vu et le second est réfléchi symétriquement par rapport à l'axe optique qui fait office de normale au miroir en S.

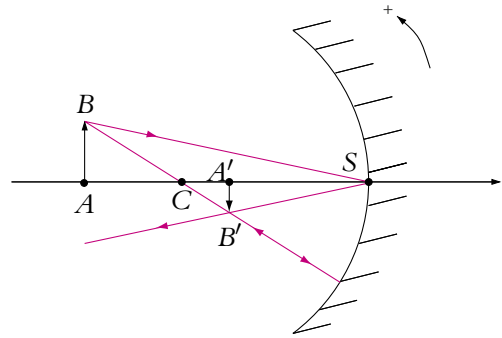


Figure 10.20 Utilisation d'un rayon passant par le sommet du miroir.

Pour calculer le grandissement, on utilise les relations de Thalès. Dans les triangles BAC et $B'A'C$, on peut écrire :

$$\gamma = \frac{\overline{A'B'}}{\overline{AB}} = \frac{\overline{CA'}}{\overline{CA}} \quad (10.8)$$

Les triangles ABS et $A'B'S'$ ont même angle en S donc en considérant le rayon qui se réfléchit en S :

$$\frac{\overline{A'B'}}{\overline{SA'}} = -\frac{\overline{AB}}{\overline{SA}}$$

soit :

$$\gamma = \frac{\overline{A'B'}}{\overline{AB}} = -\frac{\overline{SA'}}{\overline{SA}} \quad (10.9)$$

3.10 Constructions

a) Règles de construction

On revient dans ce paragraphe sur les méthodes de construction permettant de déterminer les positions des images et des rayons réfléchis dans les conditions de Gauss.



Dans le paragraphe précédent, on a utilisé un rayon passant par le sommet mais il faut éviter d'utiliser un tel rayon pour faire une construction précise : il est difficile de construire deux angles égaux sans rapporteur. On n'utilise ce rayon que pour retrouver l'expression du grandissement comme cela a été fait.

Pour une construction précise, les idées de base sont les suivantes :

1. un rayon passant par le centre n'est pas dévié ;
2. un rayon incident parallèle à l'axe (ou son prolongement) passe après réflexion par le foyer image ;
3. un rayon incident (ou son prolongement) passant par le foyer objet est, après réflexion, parallèle à l'axe ;
4. deux rayons incidents parallèles donnent des rayons réfléchis qui, eux ou leurs prolongements, se croisent dans le plan focal image ;
5. deux rayons incidents qui, eux ou leurs prolongements, se croisent dans le plan focal objet donnent des rayons réfléchis parallèles entre eux.

► **Remarques :**

- On exagère sur le schéma les dimensions perpendiculaires à l'axe pour la clarté de ces dessins si bien que les lois de Descartes ne semblent plus vérifiées.
- On utilise en général trois rayons pour effectuer les constructions. En fait, deux suffisent mais le troisième confirme la construction.

b) Construction d'une image correspondant à un objet

On utilise les trois premières règles. Sur toutes les constructions, les données sont représentées en gras alors que ce qui est obtenu par construction est représenté en trait fin. On encourage le lecteur à refaire toutes les constructions en utilisant une feuille de papier calque et en partant des mêmes données afin de pouvoir comparer les résultats. D'autre part, il est toujours plus clair, quand on peut le faire, d'utiliser des couleurs différentes pour les rayons.

Il faut bien respecter les conventions : **les objets et images virtuelles sont représentés en pointillés ainsi que les prolongements des rayons.**

c) Miroir concave

Les trois positions possibles pour l'objet sont :

1. réel avant F ;
2. réel entre F et S ;
3. virtuel après S .

Les trois constructions sont données sur les figures 10.21, 10.22 et 10.23.

Dans la première construction (Cf. figure 10.21), on a tracé :

1. un rayon passant par B et parallèle à l'axe qui est réfléchi en passant par F ,
2. un rayon passant par B et par F qui est réfléchi parallèlement à l'axe,
3. un rayon passant par B et par C qui est réfléchi sur lui-même.

► **Remarque** : Si l'objet est avant C comme sur la figure 10.21, son image est plus petite ; s'il est entre C et F , elle est plus grande, la construction étant la même.

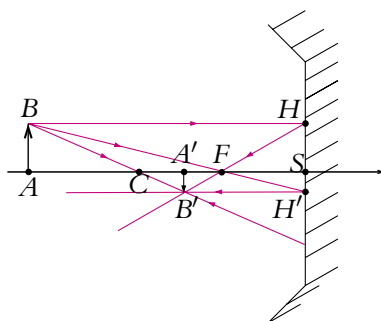


Figure 10.21 Construction de l'image d'un objet réel situé avant le foyer objet par un miroir concave.

Les trois rayons réfléchis se croisent en B' , l'image est donc réelle.

Dans la construction suivante (Cf. figure 10.22), on a tracé :

1. un rayon passant par B et parallèle à l'axe qui est réfléchi en passant par F ,
2. un rayon passant par B et par F qui est réfléchi parallèlement à l'axe,
3. un rayon passant par B et par C qui est réfléchi sur lui-même.

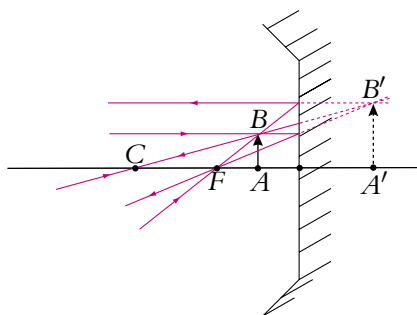


Figure 10.22 Construction de l'image d'un objet réel situé entre le foyer objet et le sommet par un miroir concave.

Les prolongements des rayons réfléchis se croisent en B' , l'image est donc virtuelle.

Dans la dernière construction (Cf. figure 10.23), on a tracé :

1. un rayon incident parallèle à l'axe, dont le prolongement arrive en B et qui est réfléchi en passant par F ,
2. un rayon incident passant par F , dont le prolongement arrive en B et qui est réfléchi parallèlement à l'axe,
3. un rayon incident passant par C , dont le prolongement arrive en B et qui est réfléchi sur lui-même.

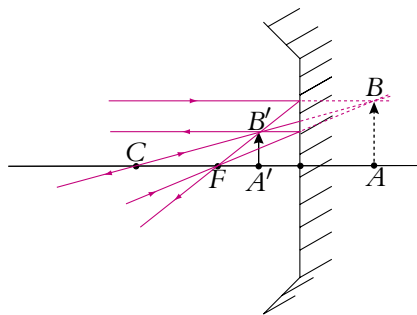


Figure 10.23 Construction de l'image d'un objet virtuel situé après le sommet par un miroir concave.

Les rayons réfléchis se croisent en B' , l'image est donc réelle.

Les résultats sont les suivants :

1. un objet réel avant F donne une image réelle renversée,
2. un objet réel entre F et S donne une image virtuelle droite plus grande, c'est le miroir grossissant,
3. un objet virtuel après S donne une image réelle plus petite droite.



Attention à la construction dans le cas d'un objet virtuel. Il faut se souvenir que cet objet est l'image formée par un autre instrument non représenté ici et que l'on a placé le miroir avant cette image devenue un objet pour lui. Ce sont donc les prolongements des rayons incidents qui doivent se croiser sur l'objet et non les rayons émergents.

d) Miroir convexe

Les trois positions possibles pour l'objet sont :

1. réel avant S ;
2. virtuel entre S et F ;
3. virtuel après F .

Les trois constructions sont données sur les figures 10.24, 10.25 et 10.26.

Dans la première construction (Cf. figure 10.24), on a tracé :

1. un rayon incident, passant par B et parallèle à l'axe, et le rayon réfléchi correspondant dont le prolongement passe par F ,
2. un rayon incident, passant par B , dont le prolongement passe par F , qui est réfléchi parallèlement à l'axe,
3. un rayon incident passant par B , dont le prolongement passe par C , qui est réfléchi sur lui-même.

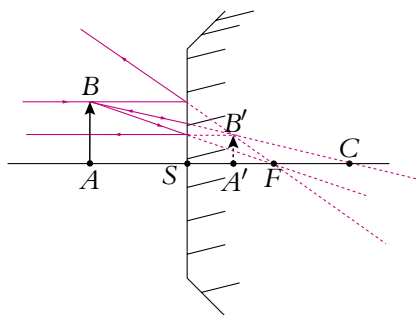


Figure 10.24 Construction de l'image d'un objet réel situé avant le sommet par un miroir convexe.

Les prolongements des rayons réfléchis se croisent en B' , l'image est donc virtuelle.

Dans la construction suivante (Cf. figure 10.25), on a tracé :

1. un rayon incident parallèle à l'axe dont le prolongement passe par B et le rayon réfléchi correspondant dont le prolongement passe par F ,
2. un rayon incident, dont le prolongement passe par B et F , qui est réfléchi parallèlement à l'axe,
3. un rayon incident, dont le prolongement passe par B et C , qui est réfléchi sur lui-même.

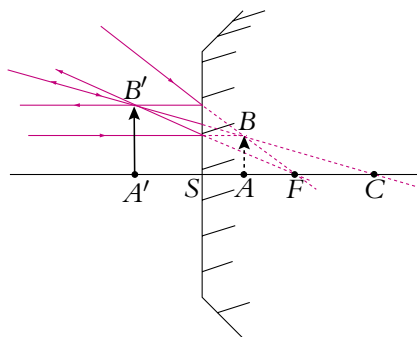


Figure 10.25 Construction de l'image d'un objet virtuel situé entre le sommet et le foyer par un miroir convexe.

Les rayons réfléchis se croisent en B' , l'image est donc réelle.

Dans la dernière construction (Cf. figure 10.26), on a tracé :

1. un rayon incident parallèle à l'axe dont le prolongement passe par B et le rayon réfléchi correspondant dont le prolongement passe par F ,
2. un rayon incident, dont le prolongement passe par B et F , qui est réfléchi parallèlement à l'axe,
3. un rayon incident, dont le prolongement passe par B et C , qui est réfléchi sur lui-même.

► **Remarque** : Si l'objet est à droite de C comme sur la figure 10.26, l'image est plus petite ; s'il est entre C et F , elle est plus grande, la construction étant la même.

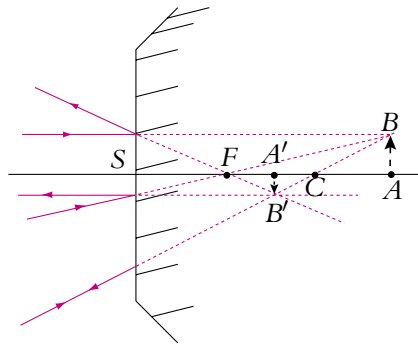


Figure 10.26 Construction de l'image d'un objet virtuel situé après le foyer objet par un miroir convexe.

Les prolongements des rayons réfléchis se croisent en B' , l'image est donc virtuelle.

Les résultats des trois constructions sont les suivants :

1. un objet réel avant S donne une image virtuelle droite plus petite,
2. un objet virtuel entre S et F donne une image réelle droite plus grande,
3. un objet virtuel après F donne une image virtuelle renversée.

Dans le chapitre sur la focométrie, on donnera un moyen rapide de reconnaître un miroir concave ou un miroir convexe.

e) Construction d'un rayon réfléchi correspondant à un incident donné

On utilise les règles 4 et 5. Les constructions ont été réalisées sur des figures différentes pour plus de lisibilité. Le rayon incident pour lequel on cherche le rayon réfléchi est en trait gras. Dans tous les cas, on utilise le rayon passant par le centre qui est facile à tracer.

Pour le miroir concave :

1. Sur la figure 10.27, on trace le rayon parallèle au rayon incident donné qui passe par C . Ce rayon est réfléchi sur lui-même et croise le rayon réfléchi recherché au foyer secondaire image F'_s , intersection du plan focal image et de ce rayon passant par C .
2. Sur la figure 10.28, on trace le rayon passant par C qui coupe le rayon incident donné dans le plan focal objet, au point F_s , foyer secondaire objet. Les rayons réfléchis sont parallèles.

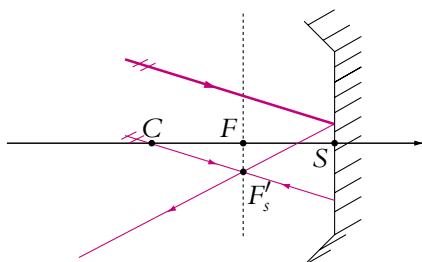


Figure 10.27 Utilisation d'un foyer image secondaire.

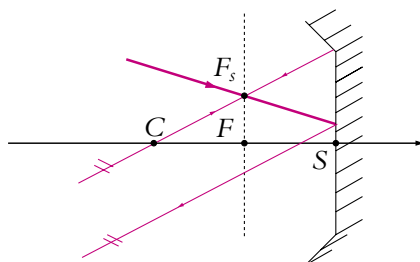


Figure 10.28 Utilisation d'un foyer objet secondaire.

Pour le miroir convexe :

1. Sur la figure 10.29, on trace le rayon parallèle au rayon incident donné et dont le prolongement passe par C . Ce rayon est réfléchi sur lui-même et croise le prolongement du rayon réfléchi recherché au foyer secondaire image F'_s , intersection du plan focal image et de ce rayon passant par C .
2. Sur la figure 10.30, on trace le rayon dont le prolongement passe par C qui coupe le prolongement du rayon incident donné dans le plan focal, au point F_s , foyer secondaire objet. Les rayons réfléchis sont parallèles.

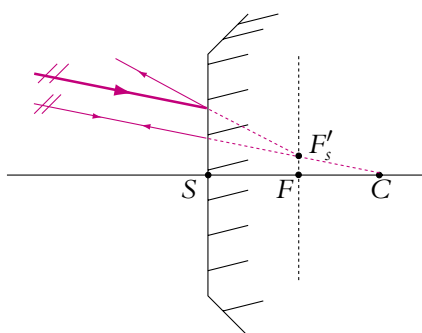


Figure 10.29 Utilisation d'un foyer image secondaire.

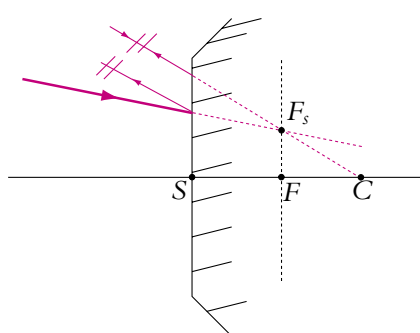


Figure 10.30 Utilisation d'un foyer objet secondaire.

3.11 Complément : relation de Lagrange-Helmholtz

Il s'agit d'établir une relation entre les tailles des objets et des images et les angles que font les rayons avec l'axe. On utilise pour cela les notations de la figure 10.31.

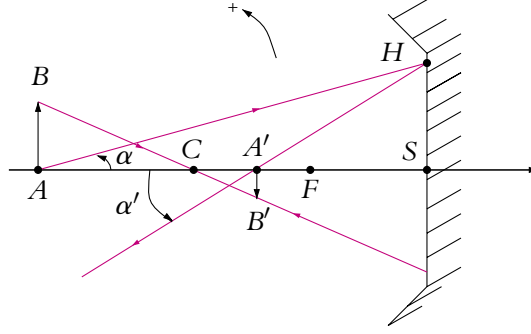


Figure 10.31 Relation de Lagrange-Helmholtz d'un miroir.

Alors :

$$\tan \alpha = \frac{\overline{SH}}{\overline{AS}} \quad \text{et} \quad \tan \alpha' = \frac{\overline{SH}}{\overline{A'S}}$$

D'autre part, dans les conditions de Gauss, les angles sont petits donc $\tan \alpha \simeq \alpha$ et $\tan \alpha' \simeq \alpha'$. Les relations précédentes donnent :

$$\frac{\alpha'}{\alpha} = \frac{\overline{AS}}{\overline{A'S}}$$

Or la relation de grandissement (10.9) donne :

$$\frac{\overline{AS}}{\overline{A'S}} = -\frac{\overline{AB}}{\overline{A'B'}}$$

En combinant les relations précédentes, on obtient la relation de Lagrange-Helmholtz :

$$\alpha \overline{AB} = -\alpha' \overline{A'B'} \quad (10.10)$$

On définit le *grossissement* (ou grandissement angulaire) par $G = \frac{\alpha'}{\alpha}$. On peut déduire de la relation de Lagrange-Helmholtz une relation entre le grandissement et le grossissement :

$$G = \frac{\alpha'}{\alpha} = -\frac{\overline{AB}}{\overline{A'B'}} = -\frac{1}{\gamma} \quad (10.11)$$

ou encore :

$$G \gamma = -1 \quad (10.12)$$

A. Applications directes du cours

1. Images réelles avec un miroir

1. Déterminer les positions des objets ayant une image réelle par un miroir sphérique concave. Réaliser les constructions correspondantes.
2. Déterminer les positions des objets ayant une image réelle par un miroir sphérique convexe. Réaliser les constructions correspondantes.

2. Caractérisation d'un miroir

Un miroir sphérique donne d'un objet réel situé à 60 cm de son sommet une image droite et réduite d'un facteur 5. En déduire, par le calcul et par une construction géométrique qu'on expliquera, les caractéristiques du miroir.

3. Question à choix multiples - Étude d'un miroir sphérique, d'après ENAC Pilotes 2004

Un miroir sphérique de centre C et de sommet S est plongé dans un milieu homogène et isotrope d'indice n . Les distances algébriques sont comptées positivement dans le sens de propagation de la lumière incidente.

1. Donner les positions des foyers objet F et image F' du miroir :
 1. F est au milieu du segment SC et F' est symétrique de F par rapport au sommet S ,
 2. F' est au milieu du segment SC et F est symétrique de F' par rapport au centre C ,
 3. F et F' sont confondus et situés au milieu du segment SC ,
 4. F et F' sont rejetés à l'infini.
2. Quelle doit être la vergence $V = \frac{1}{SF}$ (exprimée en dioptries, de symbole δ , avec $1 \delta = 1 \text{ m}^{-1}$) d'un miroir sphérique placé dans l'air pour qu'il donne d'un objet réel situé à 10 m du sommet une image droite et réduite d'un rapport 5 ?
 1. $V = 0,4 \delta$,
 2. $V = -12 \delta$,
 3. $V = 3,7 \delta$,
 4. $V = 12 \delta$.
3. Quelle est la nature d'un tel miroir ?
 1. convergent et convexe,
 2. divergent et concave,
 3. divergent et convexe,
 4. convergent et concave.
4. Un objet est placé dans un plan orthogonal à l'axe optique du miroir passant par le centre C . Où se trouve l'image ?
 1. dans le même plan passant par le centre C du miroir,

2. dans le plan focal image du miroir,
3. à l'infini,
4. dans le plan passant par le sommet S du miroir.
5. Exprimer dans ce cas le grandissement :
 1. $G = 1$,
 2. $G = -\frac{1}{2}$,
 3. $G = 2$,
 4. $G = -1$.

4. Comment se voir en entier dans un miroir

On considère un miroir plan de hauteur h accroché à un mur vertical. Une personne de taille t a ses yeux à une hauteur o du sol et se trouve à une distance d du miroir. On suppose que le bord supérieur du miroir est à la hauteur $\frac{t+o}{2}$, c'est-à-dire à mi-hauteur entre le sommet de la tête de la personne et de ses yeux.

1. a) Déterminer graphiquement la partie de son corps que la personne voit d'elle-même dans le miroir.
- b) Est-ce que la partie visible du corps dans le miroir dépend de la distance d ?
2. Quelle est la hauteur minimale h du miroir permettant à une personne de se voir entièrement ?
3. Pourquoi a-t-on accroché la bord supérieur du miroir à mi-distance entre le sommet de la tête et les yeux ?

5. Miroirs domestiques

Les miroirs d'utilisation courante sont en général des lames à faces parallèles en verre, d'épaisseur e et d'indice n . On métallise la face arrière de la lame de manière à réfléchir la lumière.

1. À quel système optique le miroir domestique est-il équivalent ?
2. a) Déterminer graphiquement la position du miroir plan équivalent M_e .
- b) Déterminer (dans l'approximation de Gauss), la distance de M_e à la première face de la lame.
3. Retrouver le résultat de la question précédente par l'utilisation de la formule de conjugaison du dioptré plan (dans les conditions de Gauss).

B. Exercices et problèmes

1. Le sextant

Le sextant est un appareil qui était utilisé par les marins (avant l'avènement du système G.P.S.) pour déterminer la latitude. Pour cela il mesurait la hauteur d'un astre (Soleil, Lune...) au-dessus de l'horizon.

Le sextant est composé de deux miroirs plans M_1 et M_2 (ce dernier étant semi-réfléchissant), et d'une lunette de visée L_v (l'exercice n'exige aucune connaissance sur la lunette). La direction OC du plan du miroir M_1 pointe sur un repère de l'arc de cercle gradué AB qui mesure 60° .

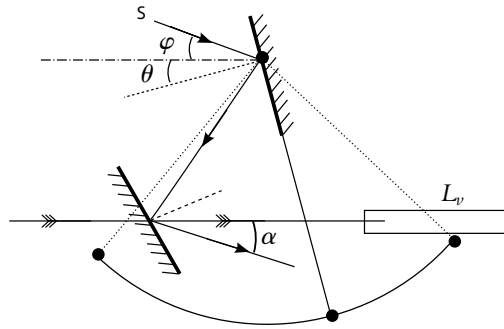


Figure 10.32

Avec la lunette, l'observateur vise l'horizon donc l'axe optique de la lunette a la direction de l'horizontale. Le plan du miroir M_2 (parallèle au segment OA) fait un angle de 30° avec cet axe. On note θ l'angle de la normale à M_1 avec l'horizontale et ϕ celui de la direction de l'astre observé S avec cette même horizontale (figure 10.32).

1. Déterminer en fonction de ϕ et θ l'angle α que fait le rayon réfléchi sur M_2 avec l'horizontale.
2. a) Quel doit-être l'angle θ pour que le rayon réfléchi sur M_2 soit selon l'axe optique de la lunette ?
b) En déduire alors la relation entre l'angle AOC et l'angle SOH .

2. Télescope à deux miroirs concaves

On considère le télescope constitué de deux miroirs sphériques concaves (figure 10.33). Les rayons provenant de l'astre observé se réfléchissent sur le miroir M_1 (sommet S_1 , centre C_1) puis sur le miroir M_2 (sommet S_2 , centre C_2).

On pose $\overline{S_2S_1} = D$, $\overline{S_1C_1} = R_1 (< 0)$ et $\overline{S_2C_2} = R_2 (> 0)$.

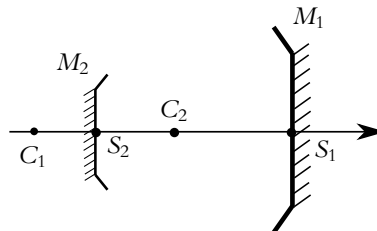


Figure 10.33

1. On nomme F le foyer objet et F' le foyer image. Déterminer les positions $\overline{S_1F}$ et $\overline{S_2F'}$ de ces deux points en fonction de D , R_1 et R_2 .
2. a) À quelle condition le système est-il afocal ?
b) Déterminer D . Application numérique pour $R_1 = -2$ m et $R_2 = 1$ m.
c) Faire une figure en respectant les rapports entre les différentes grandeurs. Construire la direction α' du faisceau émergent correspondant à un faisceau incident de direction α . En déduire le grandissement angulaire $G = \frac{\alpha'}{\alpha}$.
3. On considère la cas où $R_1 = -R_2$ et $D = -2R_1$.
a) Déterminer la position et la taille de l'image définitive d'un astre vu l'angle $2\alpha = 2^\circ$ symétriquement de part et d'autre de l'axe optique.
b) Application numérique : $R_1 = -2$ m.
c) Faire la figure en respectant les proportions et en faisant apparaître l'image intermédiaire.

3. Un élément de sécurité : le rétroviseur, d'après ENSTIM 2006

Le champ d'un miroir est la portion de l'espace qu'un observateur voit dans un miroir. Ainsi, un rétroviseur de voiture ne permet pas au conducteur de voir une autre voiture qui se situerait hors de cette portion ; c'est ce qu'on appelle l'angle mort.

1. Le rétroviseur est un miroir plan

Le rétroviseur est un miroir plan de largeur L . L'observateur place son œil, supposé ponctuel, en un point A' de l'axe du miroir, à une distance D de celui-ci.

- a) Sur une figure, positionner le point A dont l'image est A' par le miroir.
- b) Où se situent les points que l'observateur peut espérer voir par réflexion dans le miroir ? Faire apparaître cette portion de l'espace sur une construction.
- c) Préciser la valeur de l'angle α qui caractérise la portion de l'espace accessible à la vision ; c'est le champ du miroir.
- d) Application numérique : calculer α , avec $L = 20$ cm, $D = 50$ cm.

2. Le rétroviseur est un miroir sphérique

Le miroir plan est remplacé par un miroir sphérique convexe, de rayon de courbure $R = 50$ cm et de même largeur L . L'œil de l'observateur est toujours placé en A' .

- a) Effectuer la construction graphique du point A dont l'image est A' par le miroir et déterminer la position de A par rapport à S sommet du miroir.
- b) Faire apparaître le champ du miroir sur la construction.
- c) Préciser la valeur de l'angle α' qui caractérise le champ de vision.
- d) Application numérique : calculer α' avec $L = 20$ cm, $D = 50$ cm.

3. Comparaison des deux dispositifs

- a) Comparer les champs angulaires des deux types de rétroviseur.

b) Un objet, de taille 1 m, est situé à une distance $D' = 10$ m du rétroviseur. Faire une construction graphique de l'image dans les deux cas. Déterminer puis calculer les angles apparents sous lesquels l'automobiliste voit l'objet avec les deux types de rétroviseur. Commenter.

4. Télescope de Cassegrain, d'après ENSTIM 2005

- Définir le stigmatisme rigoureux d'un système optique centré.
- Qu'appelle-t-on stigmatisme approché ? À quelle(s) condition(s) est-il vérifié ?
- Définir l'aplanétisme d'un système optique.
- On considère un miroir sphérique concave de centre C et de sommet S . Un objet AB assimilable à un segment est placé perpendiculairement à l'axe optique, l'extrémité A étant située sur cet axe. Construire, dans le cadre de l'approximation de Gauss, l'image $A'B'$ de AB . La construction s'effectuera à l'aide de deux rayons émis par B l'un passant par le centre C , l'autre par le sommet S et on justifiera la trajectoire de chacun. On prendra $\overline{CS}$ égal à 5 cm et $\overline{AC}$ à 4 cm.
- Établir à partir de cette construction la formule de conjugaison des miroirs sphériques dans l'approximation de Gauss avec origine au sommet.
- En déduire l'existence d'un foyer objet F et d'un foyer image F' dont on précisera les positions relatives par rapport au sommet S et au centre C .
- On considère désormais le télescope de Cassegrain constitué de deux miroirs sphériques M_1 et M_2 . Le miroir M_1 est concave avec une ouverture à son sommet S_1 tandis que le miroir M_2 est convexe, sa face réfléchissante étant tournée vers celle de M_1 .
On observe à travers ce télescope un objet $\overline{AB}$ dont l'extrémité A est sur l'axe optique. L'objet étant très éloigné, les rayons issus de B qui atteignent le miroir M_1 sont quasiment parallèles et forment avec l'axe optique une angle α . Après réflexion sur M_1 , ces rayons se réfléchissent sur M_2 et forment une image finale $\overline{A'B'}$ située au voisinage de S_1 .
Effectuer les constructions géométriques de l'image intermédiaire $\overline{A_1B_1}$ de $\overline{AB}$ par M_1 et de l'image finale $\overline{A'B'}$. On prendra 6 cm pour $\overline{S_2S_1}$, 9 cm pour $\overline{F_1S_1}$ et 6 cm pour $\overline{F_2S_2}$ en appelant F_1 , F_2 , S_1 et S_2 respectivement le foyer de M_1 , le foyer de M_2 , le sommet de M_1 et le sommet de M_2 .
- On note $f_1 = \overline{F_1S_1}$ et $f_2 = \overline{F_2S_2}$ les distances focales comptées positivement des miroirs M_1 et M_2 . On appelle $D = \overline{S_2S_1}$ la distance séparant les deux miroirs. Exprimer D en fonction de f_1 et f_2 pour que l'image finale $\overline{A'B'}$ soit située dans le plan de S_1 .
- Que devient cette expression quand $f_1 \gg f_2$?
- Déterminer dans ces conditions (sans l'approximation $f_1 \gg f_2$) la taille de l'image intermédiaire $\overline{A_1B_1}$ en fonction de α et de f_1 .
- En déduire celle de $\overline{A'B'}$ en fonction de α , f_1 et f_2 .
- Simplifier cette expression quand $f_1 \gg f_2$.
- Déterminer les valeurs de $\overline{A_1B_1}$ et de $\overline{A'B'}$ pour $\alpha = 1,0 \cdot 10^{-3}$ rad, $f_1 = 40$ cm et $\frac{f_1}{f_2} = 20$.

Lentilles minces sphériques

11

1. Présentation

1.1 Les différentes lentilles

Une *lentille sphérique* est un milieu transparent, homogène et isotrope limité par deux dioptries sphériques ou un dioptre sphérique et un dioptre plan.

Comme les dioptries sphériques, les lentilles possèdent un axe de symétrie appelé *axe optique*. Il s'agit de la droite passant par les sommets des deux dioptries dans le cas où aucun dioptre n'est plan. Dans le cas contraire, c'est l'axe de symétrie du dioptre non plan (passant par son centre et son sommet).

On distingue deux types de lentilles :

- les lentilles à bords minces,
- les lentilles à bords épais.

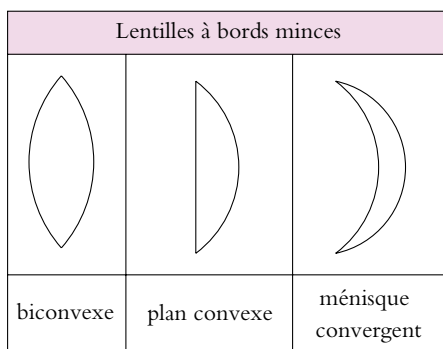


Figure 11.1 Différentes lentilles à bords minces.

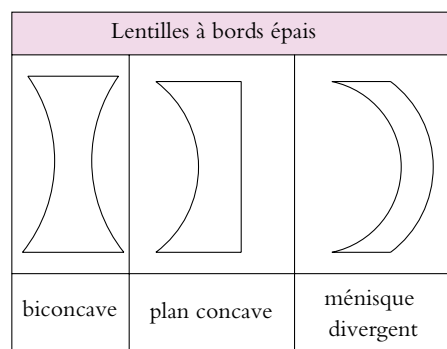


Figure 11.2 Différentes lentilles à bords épais.

Les lentilles à bords minces sont convergentes c'est-à-dire qu'elles rabattent les rayons vers l'axe alors que les lentilles à bords épais sont divergentes : elles écartent les rayons de l'axe. Les figures 11.1 et 11.2 récapitulent les six formes de lentilles.

On définit le *centre optique* (en général appelé O) comme le point de l'axe optique par lequel passe le rayon réfracté correspondant à **un rayon incident dont le rayon émergent correspondant lui est parallèle** (Cf. figure 11.3).

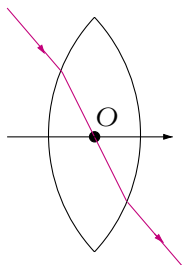


Figure 11.3 Déviation d'un rayon par une lentille.

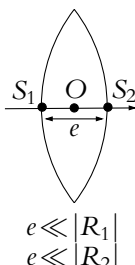
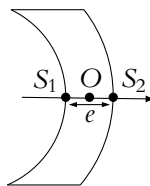
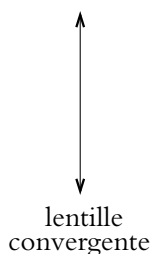


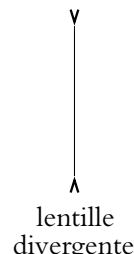
Figure 11.4 Conditions pour avoir une lentille mince.



$$e \ll |R_2 - R_1|$$



lentille convergente



lentille divergente

Figure 11.5 Symbole des lentilles minces.

Une lentille sphérique mince est une lentille sphérique dont l'épaisseur e sur l'axe est petite devant les rayons des dioptries et devant la différence de ceux-ci s'ils sont de même sens (Cf. figure 11.4). Ainsi si l'on considère les rayons algébriques R_1 et R_2 des dioptries, une lentille est dite mince si :

$$e \ll |R_1| \quad , \quad e \ll |R_2| \quad \text{et} \quad e \ll |R_1 - R_2| \quad (11.1)$$

Dans ce cas, on peut considérer que les trois points O , S_1 et S_2 sont confondus si bien qu'on représente les lentilles par un trait sur lequel on fait seulement figurer O . Les symboles des lentilles convergente et divergente sont donnés sur la figure 11.5.

1.2 Distance focale, vergence

Dans l'approximation de Gauss, les lentilles sont des systèmes optiques stigmatiques (stigmatisme approché) et aplanétiques (aplanétisme approché). On peut donc définir deux foyers, objet et image, et deux plans focaux, objet et image. Dans toute la suite de l'étude, on suppose que les lentilles sont utilisées dans l'air d'indice $n = 1$.

a) Cas d'une lentille convergente

Si on éclaire une lentille convergente par un faisceau parallèle à l'axe, on observe que les rayons réfractés convergent en un point de l'espace image qui est donc le foyer image F' (Cf. figure 11.7). De même, si on place une source ponctuelle au point F , symétrique de F' par rapport à O , on observe qu'après la lentille, le faisceau est parallèle à l'axe, donc F est le foyer objet (Cf. figure 11.6).

Pour une lentille convergente, les foyers objet et image sont symétriques par rapport au centre optique. Le foyer objet est situé avant le centre optique O et le foyer image après O dans le sens de propagation de la lumière.

Si le sens de propagation de la lumière change, les foyers sont inversés.

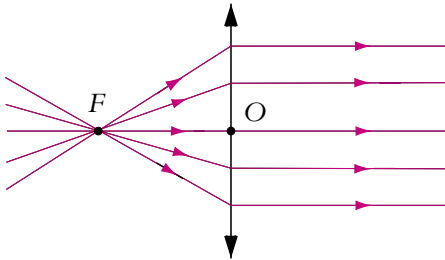


Figure 11.6 Foyer objet d'une lentille convergente.

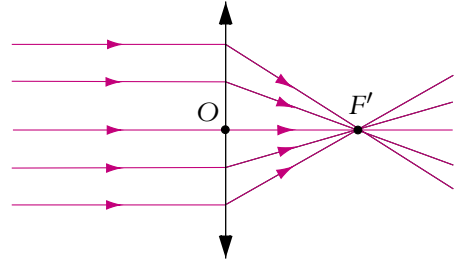


Figure 11.7 Foyer image d'une lentille convergente.

On définit les distances focales :

- la distance focale objet $f = \overline{OF}$, qui est négative pour une lentille convergente,
- la distance focale image $f' = \overline{OF'}$, qui est positive pour une lentille convergente.

Les distances focales f et f' sont opposées.

On définit la *vergence* V d'une lentille par :

$$V = \frac{1}{f'} = -\frac{1}{f} \quad (11.2)$$

La vergence est homogène à l'inverse d'une distance et l'unité de la vergence est la *dioptrie* (symbole δ) qui correspond à des m^{-1} . Ainsi une lentille de distance focale image $f' = 50 \text{ cm}$ a une vergence de 2δ .

b) Cas d'une lentille divergente

Si on éclaire une lentille divergente avec un faisceau parallèle à l'axe, on observe que le faisceau diverge après passage à travers la lentille (Cf. figure 11.9). Si on prolonge les rayons du faisceau divergent, ils semblent provenir d'un point de l'espace objet qui est donc l'image du faisceau parallèle. Ce point est le foyer image F' de la lentille, il est virtuel.

Si l'on éclaire la lentille par un faisceau convergent (Cf. figure 11.8), dont les prolongements de rayon se croisent au point symétrique de F' par rapport à la lentille, le faisceau émergent est parallèle à l'axe. Ce point est donc le foyer objet F , il est lui aussi virtuel.

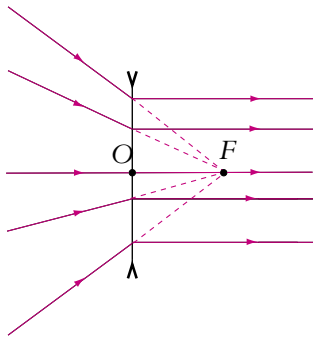


Figure 11.8 Foyer objet d'une lentille divergente.

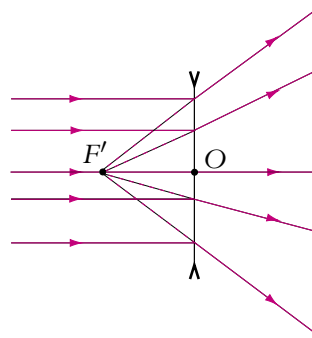


Figure 11.9 Foyer image d'une lentille divergente.

Pour une lentille divergente les foyers objet et image sont symétriques par rapport au centre optique mais virtuels. Le foyer objet est situé après le centre optique O et le foyer image avant O dans le sens de propagation de la lumière.

Si le sens de propagation de la lumière change, les foyers sont inversés comme pour une lentille convergente.

La définition de la vergence est la même que pour une lentille convergente. Dans ce cas, la distance focale image $f' = \overline{OF'}$ est négative donc la vergence $V = \frac{1}{f'}$ l'est aussi. La distance focale objet $f = -f'$ est positive.

Sur les figures précédentes, on peut vérifier la validité des termes lentilles convergentes et divergentes.

2. Relations de conjugaison

À partir des propriétés énoncées dans les paragraphes précédents, on va établir les relations de conjugaison des lentilles minces, permettant de déterminer les positions des objets et des images ainsi que les grandissements. On procède comme pour les miroirs sphériques.

Ces relations sont établies dans le cas particulier d'une lentille convergente avec un objet placé avant le foyer objet mais sont valables quelle que soit la position de l'objet ainsi que pour une lentille divergente car elles sont algébriques.

On se sert de trois règles de construction qui seront revues plus en détail dans la suite du chapitre :

- un rayon passant par le centre optique n'est pas dévié ;
- un rayon incident parallèle à l'axe émerge après la lentille en passant par le foyer image ;

- un rayon incident passant par le foyer objet émerge parallèlement à l'axe.

Ces trois règles permettent d'effectuer la construction de l'image $A'B'$ de l'objet AB sur le figure 11.10.

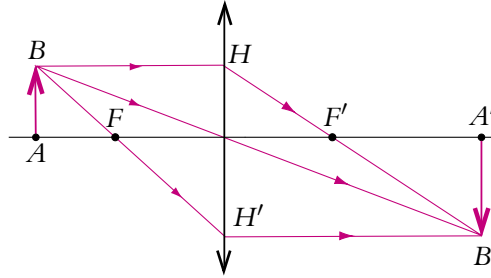


Figure 11.10 Construction d'une image par une lentille convergente.

► **Remarque :** Pour obtenir des formules faisant intervenir le centre optique, il faut tracer un rayon passant par O . De même, pour retrouver des formules faisant intervenir les foyers, il faut tracer les rayons passant par F et F' .

2.1 Relations avec origine aux foyers, dites relations de Newton

On utilise la même méthode que pour les miroirs sphériques : on calcule le grandissement en faisant intervenir le point H projeté de B sur la lentille et le point H' projeté de B' sur la lentille (Cf. figure 11.10). On peut alors calculer le grandissement de deux manières différentes. Les relations de Thalès dans les triangles ABF et $OH'F$ donnent :

$$\gamma = \frac{\overline{A'B'}}{\overline{AB}} = \frac{\overline{OH'}}{\overline{AB}} = \frac{\overline{FO}}{\overline{FA}}$$

De même, dans les triangles HOF' et $F'A'B'$:

$$\gamma = \frac{\overline{A'B'}}{\overline{AB}} = \frac{\overline{A'B'}}{\overline{OH}} = \frac{\overline{F'A'}}{\overline{F'O}}$$

En combinant les deux relations précédentes, on obtient la relation de conjugaison avec origine aux foyers appelée relation de Newton :

$$\overline{F'A'} \overline{FA} = \overline{F'O} \overline{FO} = -f'^2 = -f^2 = ff' \quad (11.3)$$

D'après les expressions précédentes et en utilisant les distances focales, on peut écrire le grandissement sous les deux formes suivantes :

$$\gamma = -\frac{\overline{F'A'}}{f'} = -\frac{f}{\overline{FA}} \quad (11.4)$$

2.2 Relations avec origine au centre, dites relations de Descartes

Les notations font encore référence à la figure 11.10. Pour obtenir le grandissement $\gamma = \frac{\overline{A'B'}}{\overline{AB}}$, il suffit d'utiliser la relation de Thalès dans les triangles OAB et $OA'B'$. On trouve alors :

$$\gamma = \frac{\overline{A'B'}}{\overline{AB}} = \frac{\overline{OA'}}{\overline{OA}} \quad (11.5)$$

Pour obtenir la relation avec origine au centre, on repart de la relation de Newton (11.3) et on introduit le point désiré O . Cela donne :

$$(\overline{F'O} + \overline{OA'}) (\overline{FO} + \overline{OA}) = -f'^2$$

En remplaçant $\overline{F'O}$ par $-f'$ et $\overline{FO}$ par $-f = f'$ et en développant, les termes en f'^2 se simplifient et on obtient :

$$f' \overline{OA'} - f' \overline{OA} + \overline{OA} \overline{OA'} = 0$$

Finalement on divise l'expression par $f' \overline{OA} \overline{OA'}$, pour trouver :

$$\frac{1}{\overline{OA'}} - \frac{1}{\overline{OA}} = \frac{1}{f'} = V \quad (11.6)$$

2.3 Complément : Relation de Lagrange - Helmholtz

Comme cela a été fait pour le miroir sphérique, on peut établir une relation entre le grossissement (rapport des angles avec et sans lentilles) et le grandissement. On trace un rayon issu de A qui fait un angle α avec l'axe optique auquel correspond un rayon faisant un angle α' avec l'axe et passant par A' puisque A' est l'image de A par la lentille (Cf. figure 11.11).

En prenant garde aux signes des angles et des distances algébriques, on écrit :

$$\tan \alpha = -\frac{\overline{OH}}{\overline{OA}} \quad \text{et} \quad \tan \alpha' = -\frac{\overline{OH}}{\overline{OA'}}$$

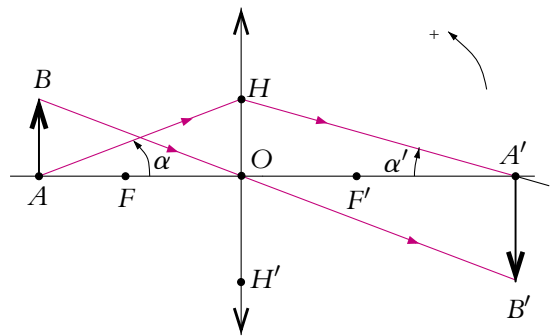


Figure 11.11 Relation de Lagrange-Helmholtz d'une lentille.

Dans les conditions de Gauss : $\tan \alpha \simeq \alpha$ et $\tan \alpha' \simeq \alpha'$ d'où la relation de Lagrange-Helmholtz :

$$\alpha \overline{OA} = \alpha' \overline{OA'} \quad (11.7)$$

Sachant que le grossissement G est défini par $G = \frac{\alpha'}{\alpha}$ et que le grandissement est $\gamma = \frac{\overline{A'B'}}{\overline{AB}} = \frac{\overline{OA'}}{\overline{OA}}$, la relation de Lagrange-Helmholtz permet d'écrire :

$$G = \frac{1}{\gamma} \quad (11.8)$$

ou encore :

$$G \gamma = 1 \quad (11.9)$$

3. Construction

Pour résoudre un problème d'optique géométrique, il est nécessaire de maîtriser parfaitement la construction des rayons lumineux. Comme cela a été fait pour les miroirs, on va récapituler tous les cas possibles de construction d'image à partir d'objet et de rayon émergent correspondant à un rayon incident.

3.1 Les règles de construction

On récapitule ici les règles de construction qui sont similaires à celles des miroirs :

1. un rayon passant par le centre optique n'est pas dévié ;
2. un rayon incident parallèle à l'axe donne un rayon émergent (ou son prolongement) qui passe par le foyer image ;
3. un rayon incident (ou son prolongement) passant par le foyer objet donne un rayon émergent parallèle à l'axe ;
4. deux rayons incidents parallèles donnent des rayons émergents qui, eux ou leurs prolongements, se croisent dans le plan focal image ;
5. deux rayons incidents qui, eux ou leurs prolongements, se croisent dans le plan focal objet donnent des rayons émergents parallèles entre eux.

On aura intérêt à se réciter ces règles à chaque fois que l'on fait une construction.

3.2 Construction d'une image correspondant à un objet

On utilise les trois premières règles. Sur toutes les constructions, les données sont représentées en gras alors que ce qui est obtenu par construction est représenté en trait fin. Comme pour les miroirs, on encourage le lecteur à refaire toutes les constructions en utilisant une feuille de papier calque pour partir des mêmes données et pouvoir

comparer les résultats. D'autre part, il est toujours plus clair, quand on peut le faire, d'utiliser des couleurs différentes pour les différents rayons.

Il faut bien respecter les conventions : **les objets et images virtuelles sont représentés en pointillés ainsi que les prolongements des rayons.**

a) Lentille convergente

Les trois positions possibles pour l'objet sont :

1. réel avant F ;
2. réel entre F et O ;
3. virtuel après O .

Les trois constructions sont données sur les figures 11.12, 11.13 et 11.14.

Sur la première construction (Cf. figure 11.12), on a tracé :

1. un rayon incident passant par B et parallèle à l'axe qui donne un rayon émergent passant par F' ;
2. un rayon incident passant par B et par F qui donne un rayon émergent parallèle à l'axe ;
3. un rayon incident passant par B et par O qui n'est pas dévié.

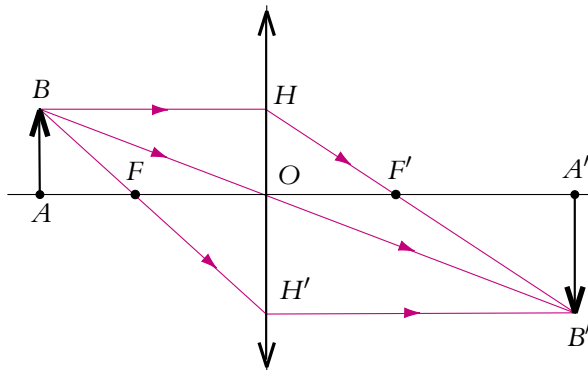


Figure 11.12 Construction de l'image d'un objet réel situé avant le foyer objet par une lentille convergente.

Les trois rayons émergents se croisent en B' , l'image est donc réelle.

Sur la construction suivante (Cf. figure 11.13), on a tracé :

1. un rayon incident passant par B et parallèle à l'axe qui donne un rayon émergent passant par F' ,
2. un rayon incident passant par B et par F , qui donne un rayon émergent parallèle à l'axe,
3. un rayon incident passant par B et par O qui n'est pas dévié.

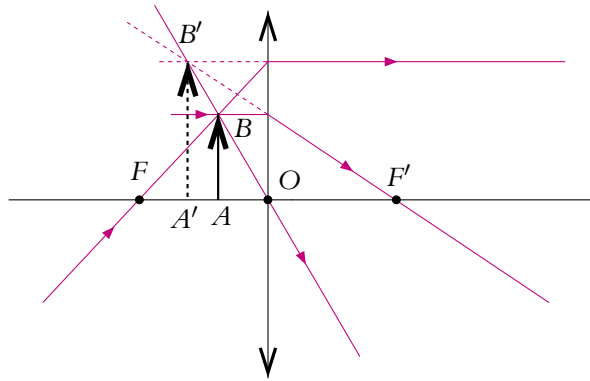


Figure 11.13 Construction de l'image d'un objet réel situé entre le foyer objet et le centre par une lentille convergente.

Les prolongements des trois rayons incidents se croisent en B' , l'image est donc virtuelle.

Sur la dernière construction (Cf. figure 11.14), on a tracé :

1. un rayon incident parallèle à l'axe dont le prolongement passe en B qui donne un rayon émergent passant par F' ,
2. un rayon incident passant par F dont le prolongement passe en B qui donne un rayon émergent parallèle à l'axe,
3. un rayon incident passant par O et B qui n'est pas dévié.

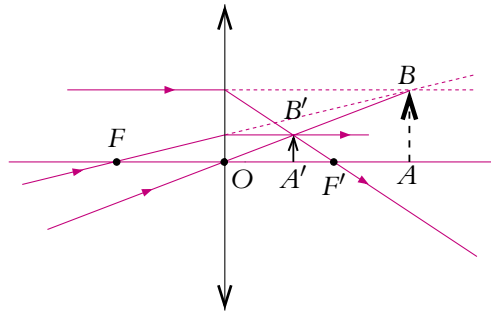


Figure 11.14 Construction de l'image d'un objet virtuel situé après le centre par une lentille convergente.

Les trois rayons émergents se croisent en B' , l'image est donc réelle.

Les résultats sont les suivants :

1. un objet réel avant F donne une image réelle, renversée, située après F' ;
2. un objet réel entre F et O donne une image virtuelle, droite, plus grande, c'est le cas de la loupe ;

3. un objet virtuel après O donne une image réelle, plus petite, droite.

Dans le cas où l'image est renversée par rapport à l'objet, le grandissement est négatif. Attention à la construction dans le cas d'un objet virtuel. Il faut se souvenir que cet objet est l'image formée par un autre instrument non représenté ici et que l'on a placé la lentille avant cette image devenue un objet pour elle. Ce sont donc les prolongements des rayons incidents qui doivent se croiser sur l'objet et non les rayons émergents.

b) Lentille divergente

Les trois positions possibles pour l'objet sont :

1. réel avant O ;
2. virtuel entre O et F ;
3. virtuel après F .

Les trois constructions sont données sur les figures 11.15, 11.16 et 11.17.

Sur la première construction (Cf. figure 11.15), on a tracé :

1. un rayon incident passant par B et parallèle à l'axe qui donne un rayon émergent dont le prolongement passe par F' ,
2. un rayon incident passant par B et dont le prolongement passe par F qui donne un rayon émergent parallèle à l'axe,
3. un rayon incident passant par B et par O qui n'est pas dévié.

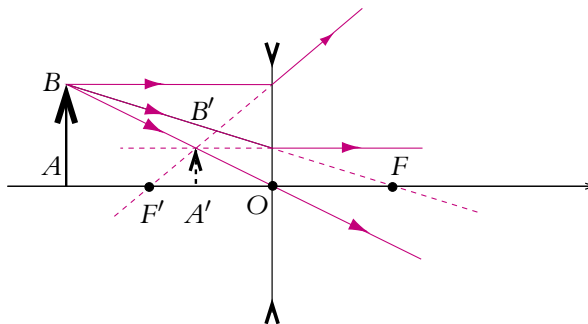


Figure 11.15 Construction de l'image d'un objet réel situé avant le centre par une lentille divergente.

Les prolongements des trois rayons incidents se croisent en B' , l'image est donc virtuelle.

Sur la construction suivante (Cf. figure 11.16) on a tracé :

1. un rayon incident parallèle à l'axe dont le prolongement passe par B qui donne un rayon émergent dont le prolongement passe par F' ,

- un rayon incident dont le prolongement passe par B et F qui donne un rayon émergent parallèle à l'axe,
- un rayon incident passant par O et B qui n'est pas dévié.

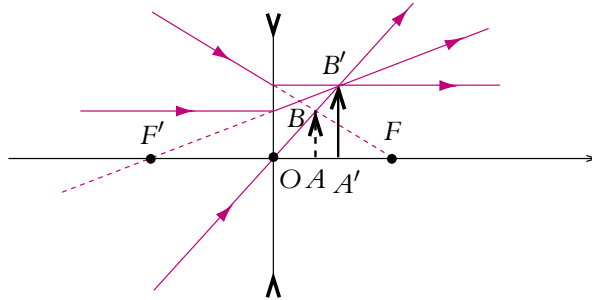


Figure 11.16 Construction de l'image d'un objet virtuel situé entre le centre et le foyer objet par une lentille divergente.

Les trois rayons émergents se croisent en B' , l'image est donc réelle.

Sur la dernière construction (Cf. figure 11.17), on a tracé :

- un rayon incident parallèle à l'axe dont le prolongement passe par B qui donne un rayon émergent dont le prolongement passe par F' ,
- un rayon incident dont le prolongement passe par B et F qui donne un rayon émergent parallèle à l'axe,
- un rayon incident passant par O et B qui n'est pas dévié.

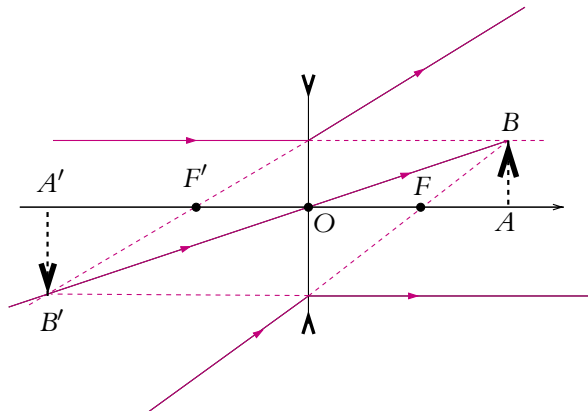


Figure 11.17 Construction de l'image d'un objet virtuel situé après le foyer objet par une lentille divergente.

Les prolongements des trois rayons émergents se croisent en B' , l'image est donc virtuelle.

Les résultats sont les suivants :

1. un objet réel avant O donne une image virtuelle, droite, plus petite ;
2. un objet virtuel entre O et F donne une image réelle, droite, plus grande ;
3. un objet virtuel après F donne une image virtuelle, renversée.

Une lentille divergente ne peut pas donner une image réelle d'un objet réel.

Pour effectuer une projection, on utilise donc une lentille convergente.

Dans le chapitre sur la focométrie, on donnera un moyen rapide de reconnaître une lentille convergente ou divergente.

► **Remarque :** Pour construire l'objet (appelé antécédent) correspondant à une image donnée, on peut soit effectuer la construction à l'envers, soit utiliser le retour inverse de la lumière en considérant l'image comme objet. Il faut alors inverser le sens de la lumière et ne pas oublier d'inverser la position des foyers.

3.3 Construction d'un rayon correspondant à un incident

Le rayon incident dont on cherche le rayon émergent est tracé en trait gras. On utilise les règles 4 et 5. Les constructions ont été réalisées sur des figures différentes pour plus de lisibilité. Dans tous les cas, on utilise le rayon passant par le centre optique qui est facile à tracer.

Pour la lentille convergente :

1. Sur la figure 11.18, on a tracé un rayon parallèle au rayon incident et passant par O . Ce rayon n'est pas dévié. Les deux rayons émergents se croisent en F'_s , intersection du plan focal image et de ce rayon passant par O .
2. Sur la figure 11.19, on a tracé un rayon incident passant par O et croisant le rayon donné au point F_s , foyer secondaire objet. Le rayon passant par O n'est pas dévié. Les deux rayons émergents sont parallèles.

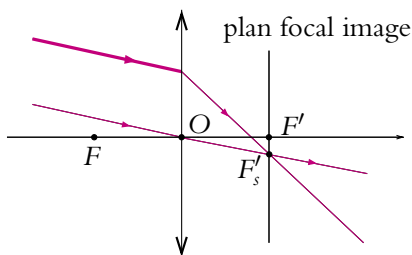


Figure 11.18 Utilisation d'un foyer image secondaire.

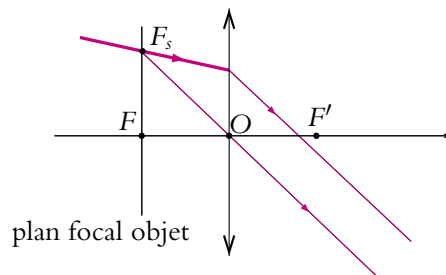


Figure 11.19 Utilisation d'un foyer objet secondaire.

Pour la lentille divergente :

1. Sur la figure 11.20, on a tracé un rayon parallèle au rayon incident et passant par O . Ce rayon n'est pas dévié. Les prolongements des deux rayons émergents se croisent en F'_s , intersection du plan focal image et de ce rayon passant par O .
2. Sur la figure 11.21, on a tracé un rayon incident passant par O et dont le prolongement croise le rayon donné au point F_s , foyer secondaire objet. Le rayon passant par O n'est pas dévié. Les deux rayons émergents sont parallèles.

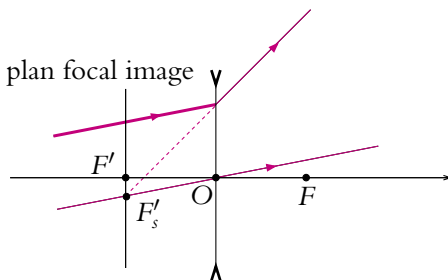


Figure 11.20 Utilisation d'un foyer image secondaire.

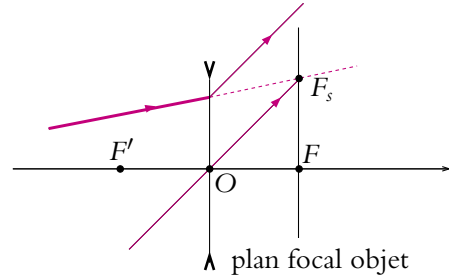


Figure 11.21 Utilisation d'un foyer objet secondaire.

3.4 Lentille accolées

Lorsqu'on accole deux lentilles minces L_1 et L_2 , de centres optiques O_1 et O_2 et de vergences V_1 et V_2 , on peut considérer que le système optique résultant est aussi une lentille mince de centre optique O confondu avec O_1 et O_2 . Quelle est la vergence V de cette lentille équivalente ?

Soit un objet A donnant une image A_1 par la première lentille (O_1 , V_1). A_1 donne une image A' par la deuxième lentille (O_2 , V_2). Si on applique les relations de conjugaison de Descartes :

$$\frac{1}{O_1 A_1} - \frac{1}{O_1 A} = V_1 \quad \text{et} \quad \frac{1}{O_2 A'} - \frac{1}{O_2 A_1} = V_2$$

Or puisque les lentilles sont minces et accolées, on peut confondre les centres optiques : $O_1 = O_2 = O$. En ajoutant les deux relations précédentes, on obtient :

$$\frac{1}{OA'} - \frac{1}{OA} = V_1 + V_2$$

On retrouve effectivement la relation de conjugaison d'une lentille mince de vergence $V = V_1 + V_2$.



Cette relation n'est valable que pour des lentilles **accollées**. Deux lentilles **non accolées** ne sont pas équivalentes à une lentille unique.

4. Quelques instruments d'optique

On va s'intéresser à trois instruments d'optique : le microscope, la lunette de Galilée et la lunette astronomique. Des développements sur la constitution des lunettes se trouvent dans le chapitre sur les instruments de travaux pratiques. On précise seulement ici qu'un instrument est constitué de deux parties optiques :

- l'*objectif* qui est du côté de l'objet ;
- l'*oculaire* qui est du côté de l'œil.

D'autre part, comme ce sera vu aussi dans le chapitre sur les instruments de travaux pratiques, pour qu'un œil observe correctement sans se fatiguer, l'objet doit être à l'infini. Or l'objet pour l'œil est l'image finale donnée par l'instrument. **Il faut donc retenir que l'image finale donnée par un instrument doit être à l'infini.**

Sur les figures, les échelles et proportions ne seront pas respectées.

4.1 Le microscope

Les valeurs numériques données sont des exemples, elles varient d'un microscope à l'autre.

On modélise le microscope par deux lentilles minces. L'objectif est équivalent à une lentille convergente L_1 de distance focale $f' = 5$ mm et l'oculaire à une lentille convergente L_2 de distance focale $f' = 2,5$ cm. La distance entre le foyer image de l'objectif F'_1 et le foyer objet de l'oculaire F_2 est notée Δ et est égale à 25 cm.

On note AB l'objet observé, A_1B_1 l'image intermédiaire donnée par l'objectif servant d'objet à l'oculaire et $A'B'$ l'image finale donnée par l'objectif (Cf. figure 11.22).

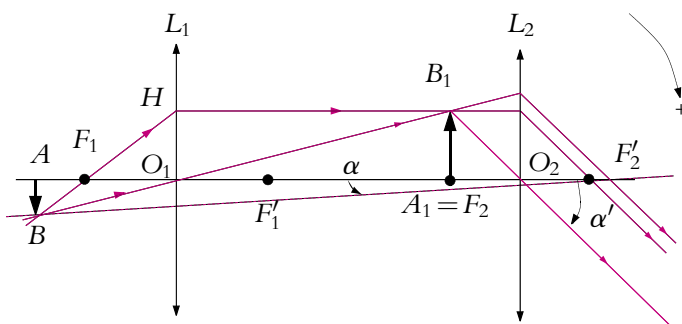


Figure 11.22 Microscope.

Puisque l'image finale est à l'infini, A_1B_1 doit se trouver dans le plan focal objet de l'oculaire :

$$AB \xrightarrow{\text{objectif}} A_1(= F_2)B_1 \xrightarrow{\text{oculaire}} \infty$$

Il s'agit tout d'abord de déterminer la position de l'objet AB par rapport à l'objectif. Puisque $A_1 = F_2$ et que l'on connaît la distance $F'_1 F_2$, il est logique d'utiliser la relation de Newton pour L_1 :

$$\overline{F'_1 F_2} \cdot \overline{F_1 A} = -f_1'^2$$

Soit :

$$\overline{F_1 A} = -\frac{f_1'^2}{\Delta} = -0,1 \text{ mm}$$

Il est logique de trouver que l'objet A est avant le foyer objet de L_1 car on veut que l'image A_1 soit réelle. Pour effectuer la construction, il faut chercher l'antécédent de F_2 par L_1 , ce qui a été réalisé sur la figure 11.22.

On obtient $\overline{O_1 A}$ par la loi de Chasles :

$$\overline{O_1 A} = \overline{O_1 F_1} + \overline{F_1 A} = -f_1' + \overline{F_1 A} = -5,1 \text{ mm}$$

Par fabrication du microscope, l'œil se place au foyer image de l'oculaire. On peut alors calculer le grossissement G qui est le rapport de l'angle α' sous lequel l'objet est vu à travers le microscope sur l'angle α sous lequel il est vu à l'œil nu. **Pour déterminer les angles sur une construction, on a intérêt à tracer les rayons passant par les centres optiques**, ce qui a été fait sur la figure 11.22. Comme on a pris le sens horaire comme sens positif pour les angles, α' est positif. On peut alors écrire :

$$\tan \alpha' \simeq \alpha' = \frac{\overline{A_1 B_1}}{\overline{F_2 O_2}} = \frac{\overline{A_1 B_1}}{f_2'}$$

Il faut alors exprimer α' en fonction de $\overline{AB}$ en utilisant les relations de grandissement pour la lentille L_1 :

$$\gamma_1 = \frac{\overline{A_1 B_1}}{\overline{AB}} = \frac{\overline{O_1 H}}{\overline{AB}} = \frac{\overline{F_1 O_1}}{\overline{F_1 A}} = \frac{f_1'}{\overline{F_1 A}}$$

D'où finalement :

$$\alpha' = \frac{\overline{AB} f_1'}{\overline{F_1 A} f_2'}$$

Il faut maintenant déterminer l'angle sous lequel AB est vu sans microscope (Cf. figure 11.23). L'angle α est négatif et :

$$\tan \alpha \simeq \alpha = \frac{\overline{AB}}{\overline{AF'_2}}$$

puisque l'œil est en F'_2 . Or

$$\overline{AF'_2} = \overline{AF_1} + \overline{F_1 F'_1} + \overline{F'_1 F_2} + \overline{F_2 F'_2} = \overline{AF_1} + 2f_1' + \Delta + 2f_2'.$$

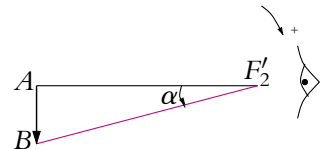


Figure 11.23 Angle sous lequel on voit un objet sans microscope.

Le grossissement est donc :

$$G = \frac{\alpha'}{\alpha} = \frac{f_1' \overline{AF_2'}}{f_2' \overline{F_1A}} = -620 \quad (11.10)$$

On trouve un grandissement négatif car l'image sera vue renversée par rapport à l'objet. En effet, sans le microscope, les rayons issus de l'objet proviennent du bas, alors qu'avec le microscope, ils semblent provenir du haut.

4.2 La lunette de Galilée

La lunette de Galilée est la lunette la plus simple qui permette d'observer des objets terrestres, c'est-à-dire sans les renverser. On modélise l'objectif par une lentille convergente L_1 de distance focale $f_1' = 60$ cm et l'oculaire par une lentille divergente L_2 de distance focale $f_2' = -5$ cm. Les valeurs ci-dessus sont des exemples, elles peuvent différer d'une lunette à l'autre.

Les objets observés sont à une distance grande devant la distance focale, on peut alors les considérer à l'infini. Comme pour le microscope, l'image finale doit être à l'infini.

L'objet et l'image étant à l'infini : on dit que la lunette est réglée à l'infini. Il s'agit d'un système afocal. Puisque l'objet est à l'infini, l'image intermédiaire est dans le plan focal image de L_1 et comme l'image finale est à l'infini, l'image intermédiaire est dans le plan focal objet de L_2 :

$$\infty \xrightarrow{\text{objectif}} A_1 (= F_2 = F_1') B_1 \xrightarrow{\text{oculaire}} \infty$$

Il faut donc disposer les lentilles de manière à ce que $F_1' = F_2$ (Cf. figure 11.24).

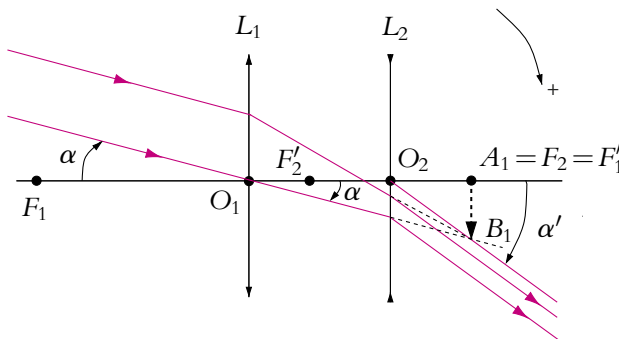


Figure 11.24 Lunette de Galilée.

Sur la figure, on a représenté un faisceau parallèle provenant de l'objet sous un angle α et émergeant de la lunette sous un angle α' . Comme cela a été fait, on a toujours intérêt à représenter l'image intermédiaire car cela peut faciliter les calculs ultérieurs.

On cherche à calculer le grossissement :

$$G = \frac{\alpha'}{\alpha}$$

On utilise pour faire ce calcul l'image intermédiaire et les rayons passant par les centres optiques :

$$\tan \alpha' \simeq \alpha' = \frac{\overline{B_1 A_1}}{\overline{O_2 F_2}} \quad \text{et} \quad \tan \alpha \simeq \alpha = \frac{\overline{B_1 A_1}}{\overline{O_1 F_1'}}$$

Avec le sens horaire choisi, $\alpha' > 0$ et $\alpha > 0$ donc :

$$G = \frac{\alpha'}{\alpha} = \frac{\overline{O_1 F_1'}}{\overline{O_2 F_2}} = \frac{f_1'}{f_2} = 12 \quad (11.11)$$

Le grossissement est positif donc objet et image sont dans le même sens.

4.3 La lunette astronomique

La lunette astronomique est aussi une lunette permettant d'observer des objets à très grandes distances (à l'infini). On modélise l'objectif par une lentille convergente L_1 de distance focale $f_1' = 60$ cm et l'oculaire par une lentille convergente L_2 de distance focale $f_2' = 5$ cm. Les valeurs ci-dessus sont des exemples, elles peuvent différer d'une lunette à l'autre.

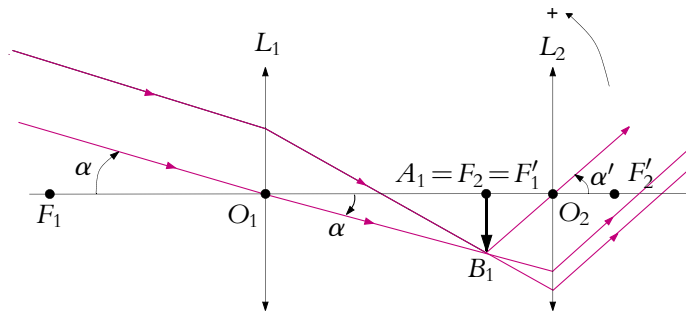


Figure 11.25 Lunette astronomique.

Comme la lunette de Galilée, la lunette astronomique est réglée à l'infini et $F_1' = F_2$ (Cf. figure 11.25) et c'est un système afocal. La méthode de calcul du grandissement est la même :

$$\tan \alpha' \simeq \alpha' = -\frac{\overline{A_1 B_1}}{\overline{F_2 O_2}} \quad \text{et} \quad \tan \alpha \simeq \alpha = \frac{\overline{A_1 B_1}}{\overline{O_1 F_1'}}$$

Avec le sens trigonométrique choisi, $\alpha' > 0$ et $\alpha < 0$ donc :

$$G = \frac{\alpha'}{\alpha} = -\frac{\overline{O_1 F_1'}}{\overline{F_2 O_2}} = -\frac{f_1'}{f_2'} = -12 \quad (11.12)$$

Le grossissement est négatif, donc l'image est à l'envers. C'est pour cela qu'on ne se sert pas de cette lunette pour observer des objets terrestres.

4.4 Conclusion

Il faut se rappeler qu'il est utile de tracer l'image intermédiaire ainsi que les rayons passant par les centres optiques.

5. Projection d'image et aberrations

On va décrire dans ce paragraphe la méthode pour projeter correctement une image, ainsi que quelques aberrations dues à l'utilisation de lentilles en dehors des conditions de Gauss.

5.1 Projection d'image

a) Montage théorique : choix de la lentille

On dispose, par exemple, d'un banc d'optique de 1,5 m et d'un objet de 5 cm × 5 cm (en fait souvent une diapositive 2,4 cm × 3,6 cm) ainsi que de deux lentilles convergentes de distance focale $f' = 10$ cm et $f' = 20$ cm. On veut obtenir un grandissement supérieur à 4 en valeur absolue : quelle lentille choisir ?

On considère une lentille convergente de distance focale image f' et de centre optique O . Soit un objet A réel. On cherche la condition pour obtenir une image A' réelle. On note $x = \overline{OA'} > 0$ et $D = \overline{AA'} > 0$. La relation de conjugaison de Descartes donne :

$$\frac{1}{\overline{OA'}} - \frac{1}{\overline{OA}} = \frac{1}{f'} \Leftrightarrow (\overline{OA} - \overline{OA'})f' = \overline{OA} \overline{OA'}$$

Or $\overline{OA} = \overline{OA'} + \overline{A'A}$ soit :

$$x^2 - xD + Df' = 0$$

On calcule le discriminant qui doit être positif puisqu'on cherche des solutions réelles :

$$\Delta = D^2 - 4f'D = D(D - 4f') \geq 0 \Rightarrow D \geq 4f' \text{ car } D > 0$$

Pour avoir une image réelle, il faut donc que $D \geq 4f'$.

Il faut donc, lors d'une projection, une distance supérieure à $4f'$ entre l'objet et l'écran pour observer une image sur l'écran.

On obtient alors deux solutions positives :

$$\overline{OA'} = \frac{D \pm \sqrt{D^2 - 4Df'}}{2}$$

Sachant que $\overline{OA} = \overline{OA'} - D$, les deux couples objets-images sont donc :

$$\overline{OA'} = \frac{D + \sqrt{D^2 - 4Df'}}{2} > 0 \quad \text{et} \quad \overline{OA} = \frac{-D + \sqrt{D^2 - 4Df'}}{2} < 0$$

et

$$\overline{OA'} = \frac{D - \sqrt{D^2 - 4Df'}}{2} > 0 \quad \text{et} \quad \overline{OA} = \frac{-D - \sqrt{D^2 - 4Df'}}{2} < 0$$

Dans les deux cas, les grandissements sont négatifs et les images renversées. Dans le premier cas, le grandissement γ_1 est tel que :

$$|\gamma_1| = \left| \frac{\overline{OA'}}{\overline{OA}} \right| = \frac{D + \sqrt{D^2 - 4Df'}}{D - \sqrt{D^2 - 4Df'}} > 1$$

et dans le second cas :

$$|\gamma_2| = \left| \frac{\overline{OA'}}{\overline{OA}} \right| = \frac{D - \sqrt{D^2 - 4Df'}}{D + \sqrt{D^2 - 4Df'}} < 1$$

Dans le cas présent, on dispose sur le banc d'une distance $D = 1,5$ m donc il faut prendre une lentille de distance focale inférieure à 37,5 cm.

On souhaite un grandissement supérieur à 4 en valeur absolue, donc :

$$\gamma = \frac{\overline{OA'}}{\overline{OA}} < -4$$

soit :

$$|\overline{OA'}| > 4|\overline{OA}|$$

sachant que $\overline{OA} > f'$ pour avoir une image réelle, on en déduit $\overline{AA'} > 5f'$ d'où $f' < 30$ cm. Les deux lentilles dont on dispose ($f' = 10$ cm et $f' = 20$ cm) vérifient les conditions précédentes, alors laquelle choisir ? De manière à se placer dans les conditions de Gauss, il faut avoir un faisceau le moins ouvert possible (petits angles d'incidence sur la lentille) : on choisit donc la lentille la moins convergente c'est-à-dire celle de focale 20 cm.

Si on réalise le montage avec les positions calculées d'après la relation de conjugaison, on s'aperçoit que l'image n'est pas très lumineuse et que l'éclairement n'est pas uniforme. Pour pallier ce problème, **il faut utiliser un condenseur**.

Un condenseur est une lentille convergente (ou un ensemble de lentilles) de faible focale permettant de concentrer les rayons de la source de lumière vers l'objet pour

qu'il soit plus lumineux (Cf. figure 11.26). Le miroir présent derrière l'ampoule sert à renvoyer vers l'avant la lumière émise vers l'arrière.

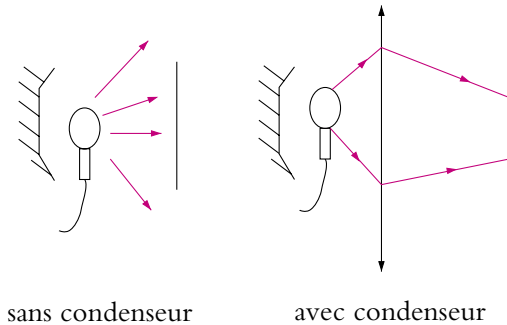


Figure 11.26 Utilisation d'un condenseur pour former une image.

b) Problème de la distorsion

Lorsqu'on veut que les images projetées soient grandes, on peut être gêné par le fait que la lentille de projection n'est plus utilisée rigoureusement dans les conditions de Gauss et que l'aplanétisme approché n'est plus vérifié. Dans ce cas, l'image projetée sur un écran ne peut être simultanément nette au centre et sur les bords. D'autre part on est souvent amené à utiliser un diaphragme pour contrôler l'éclairement. La question est de savoir où placer le diaphragme. La figure 11.27 propose trois positions dans le cas de l'image d'une grille.

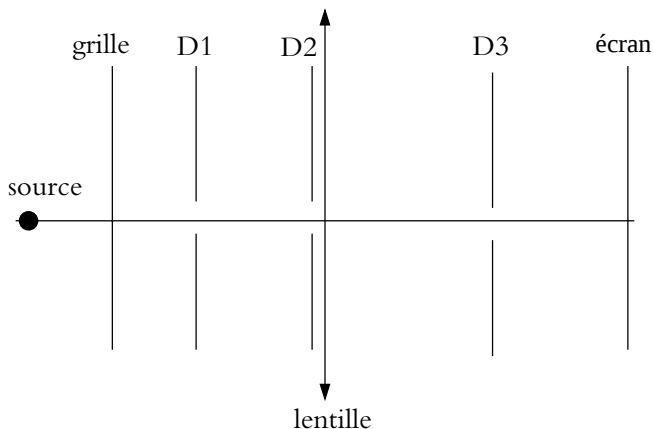


Figure 11.27 Différentes positions du diaphragme lors de la formation d'une image par une lentille convergente.

On réalise l'image d'une grille sur un écran avec une lentille de focale $f' = 12,5$ cm. On place un diaphragme aux trois positions proposées.

Dans les positions (D1) et (D2) si le diaphragme est grand ouvert, l'image est à peu près correcte bien que floue sur les bords. Par contre, si le diaphragme est très fermé, l'image est bien sûr moins lumineuse mais surtout on observe une déformation de l'image. Dans la position (D1) l'image de la grille est déformée en « barillet » (Cf. figure 11.29) alors que dans la position (D2), l'image de la grille est déformée en « coussinet » (Cf. figure 11.28). Si le diaphragme est disposé contre la lentille, la distorsion n'apparaît plus lorsqu'on le ferme.

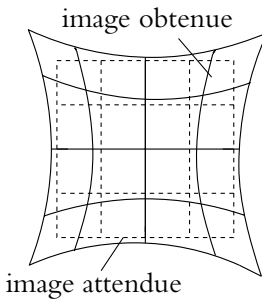


Figure 11.28 Déformation en coussinet.

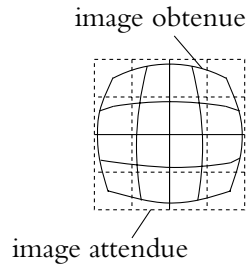


Figure 11.29 Déformation en barillet.

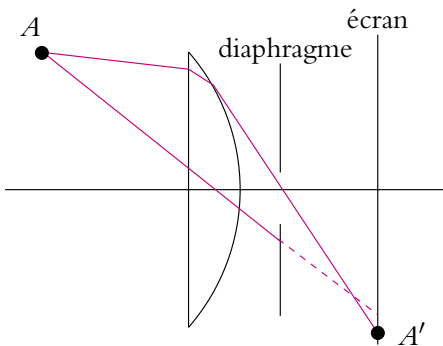


Figure 11.30 Coussinet.

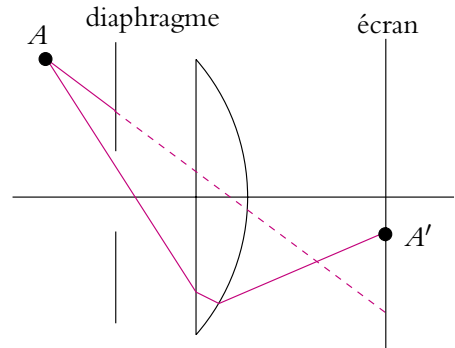


Figure 11.31 Barillet.

Ceci est dû au fait que lorsque le diaphragme est fermé, parmi les rayons qui forment la partie extérieure de l'image, on sélectionne des rayons plus obliques, c'est-à-dire qui passent plus sur les bords de la lentille. Or les bords de la lentille rabattent plus les rayons vers l'axe d'où les déformations observées sur les bords de l'image (il n'y a plus d'aplatissement car ces rayons ne vérifient pas l'approximation de Gauss). Dans le cas du barillet, le rayon extérieur arrive sur l'écran plus près de l'axe optique que le rayon passant par le centre optique (Cf. figure 11.31) et c'est le contraire dans le cas du coussinet (Cf. figure 11.30).

Cela confirme que, pour utiliser une lentille dans les conditions de Gauss, il faut que les rayons passent près du centre optique.

D'autre part, on pourra vérifier sur un schéma que, pour que l'angle entre les rayons et les faces de la lentille soit le plus petit possible, il faut placer le côté plan d'une lentille plan convexe du côté où le faisceau est le plus convergent donc placer le côté plan de la lentille du côté de l'objet ou de l'image le plus proche.

c) Montage correct pour effectuer une projection

On en déduit la méthode pour faire une projection correcte (Cf. figure 11.32) :

- faire par le condenseur l'image de la source sur la lentille ;
- tourner le côté plan de la lentille du côté du faisceau le plus convergent ;
- placer l'objet de manière à ce qu'il soit éclairé uniformément et totalement et que son image soit nette sur l'écran ;
- placer le diaphragme contre la lentille et le fermer de manière à avoir l'éclairement souhaité.

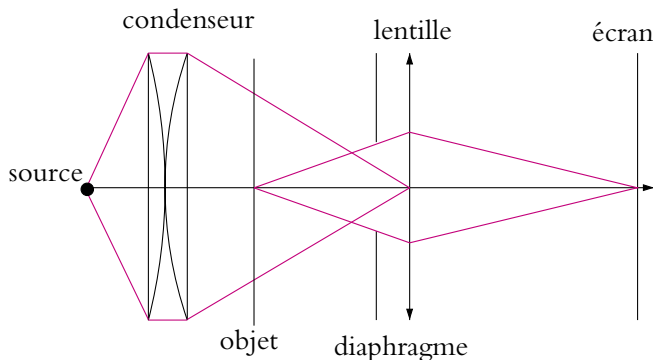


Figure 11.32 Montage pratique pour la formation d'une image.

5.2 Quelques aberrations

a) Aberrations chromatiques

Elles sont dues au fait que l'indice n du verre des lentilles dépend de la longueur d'onde λ suivant la loi de Cauchy :

$$n = a + \frac{b}{\lambda^2} \quad (11.13)$$

où a et b sont des constantes dépendant du matériau. Comme la distance focale f' est fonction de l'indice, elle dépend aussi de la longueur d'onde.

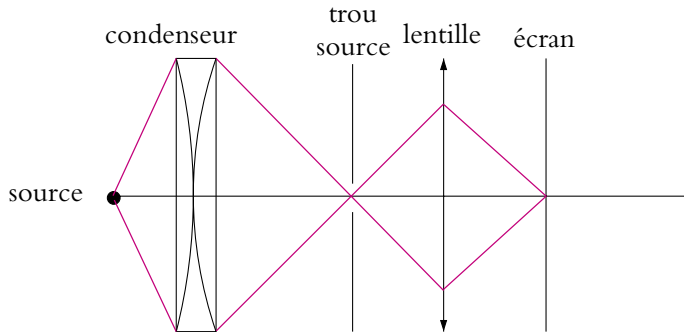


Figure 11.33 Dispersion de la lumière par une lentille : aberrations chromatiques

Si l'on effectue le montage de la figure 11.33, avec une source de lumière blanche, on observera en déplaçant l'écran de gauche à droite d'abord l'image rouge puis l'image bleue du trou source : il s'agit des aberrations longitudinales qui ont lieu le long de l'axe. Plus précisément, de gauche à droite, le faisceau sera rouge au centre et bleu autour, puis bleu au centre et rouge autour comme le montrent les figures 11.34. On parle d'aberrations latérales pour ce phénomène d'irisation. Pour éviter d'avoir des images différentes suivant la couleur, on peut utiliser des lentilles dites achromatiques qui sont formées de verre de différents indices de manière à compenser les déviations différentes suivant la couleur.

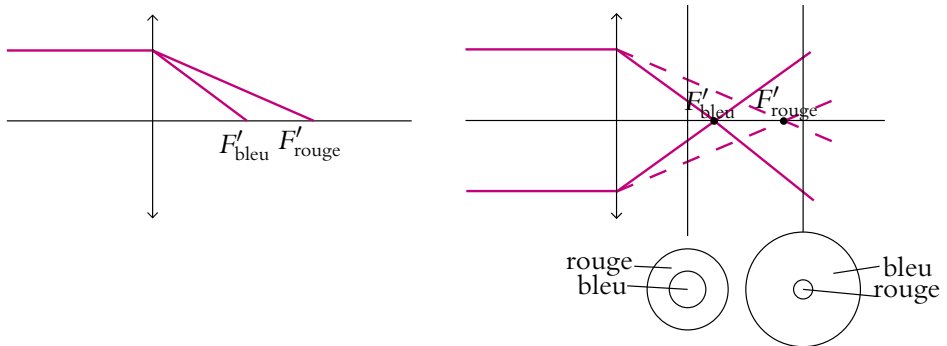


Figure 11.34 Aberrations chromatiques longitudinales et latérales d'une lentille convergente.

b) Astigmatisme et aberration de coma

À partir du montage précédent, si l'on tourne légèrement la lentille par rapport à son axe (Cf. figure 11.35), on peut observer la déformation liée à l'astigmatisme : on observe un trait horizontal au lieu de l'image ponctuelle du trou source puis en éloignant l'écran de la lentille, on observe une image ayant la forme d'un disque puis un trait vertical.

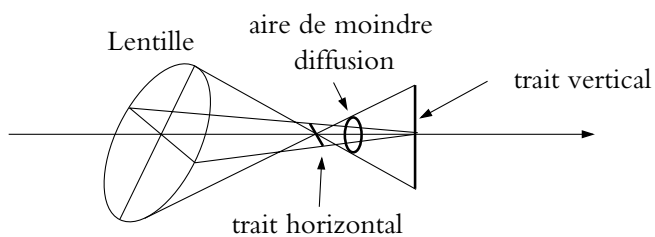


Figure 11.35 Astigmatisme et aberration de coma.

Si on tourne fortement la lentille par rapport à son axe, on observe une tâche avec une « queue » : c'est l'aberration de coma.

A. Applications directes du cours

1. Image virtuelle avec une lentille

- Déterminer les positions des objets ayant une image virtuelle par une lentille convergente.
- Déterminer les positions des objets ayant une image virtuelle par une lentille divergente.

2. Caractéristiques d'une lentille

Une lentille mince sphérique donne d'un objet réel situé à 60 cm avant son centre une image droite réduite d'un facteur 5.

Déterminer par le calcul et par une construction géométrique la position de l'image et les caractéristiques de la lentille.

3. Utilisation d'une lentille divergente

Une lentille mince divergente a pour distance focale image $f' = -30$ cm.

- Déterminer l'image d'un point A situé à 30 cm devant la lentille mince.
- Si un objet AB dans le plan transverse en A a pour taille 1 mm, quelle est la taille de l'image ?

4. Question à choix multiple - Étendeur de faisceau, d'après ENAC 2003

Un étendeur de faisceau est constitué dans cet ordre de trois lentilles minces L_1 convergente, L_2 divergente et L_3 convergente de centre O_1 , O_2 et O_3 et de distance focale f_1 , f_2 et f_3 . On note $\Delta = \overline{O_1O_2}$ et $\delta = \overline{O_2O_3}$. Ces distances Δ et δ sont telles que le système donne d'un faisceau cylindrique de rayon r_1 parallèle à l'axe optique un faisceau cylindrique de rayon $r_3 > r_1$ parallèle à l'axe optique.

- Ce système est dit :
 - divergent
 - afocal
 - convergent
 - catadioptrique
- Établir une relation entre Δ , δ , f_1 , f_2 et f_3 .
 - $\frac{1}{\delta - f_3} - \frac{1}{f_1 - \Delta} = \frac{1}{f_2}$
 - $\frac{1}{\Delta - f_2} - \frac{1}{f_1 - \delta} = \frac{1}{f_3}$
 - $\frac{1}{\delta - f_1} + \frac{1}{\Delta - f_2} = \frac{1}{f_3}$
 - $-\frac{1}{\Delta - \delta} + \frac{1}{f_3 - f_2} = \frac{1}{f_1}$
- Exprimer le rapport $\frac{r_3}{r_1}$.

$$1. \frac{r_3}{r_1} = -\frac{f_2}{f_3} \left(\frac{\delta - f_1}{f_3 - \Delta} \right)$$

$$2. \frac{r_3}{r_1} = -\frac{f_1}{f_3} \left(\frac{\Delta - f_2}{f_1 - \delta} \right)$$

$$3. \frac{r_3}{r_1} = -\frac{f_2}{f_1} \left(\frac{f_2 - f_1}{f_3 - f_1} \right)$$

$$4. \frac{r_3}{r_1} = \frac{f_3}{f_1} \left(\frac{f_1 - \Delta}{\delta - f_3} \right)$$

4. En déduire la valeur de Δ en fonction de f_1, f_2, f_3, r_1 et r_3 .

$$1. \Delta = f_1 + f_2 \left(1 - \frac{r_3 f_1}{r_1 f_3} \right)$$

$$2. \Delta = f_3 + f_1 \left(1 + \frac{r_1 f_2}{r_3 f_3} \right)$$

$$3. \Delta = f_1 + f_3 \left(1 + \frac{r_1 f_3}{r_3 f_1} \right)$$

$$4. \Delta = f_3 + f_2 \left(1 + \frac{r_3 f_2}{r_1 f_1} \right)$$

5. Même question pour δ .

$$1. \delta = f_3 + f_1 \left(1 + \frac{r_3 f_2}{r_1 f_3} \right)$$

$$2. \delta = f_2 + f_3 \left(1 + \frac{r_1 f_1}{r_3 f_3} \right)$$

$$3. \delta = f_3 + f_2 \left(1 - \frac{r_1 f_3}{r_3 f_1} \right)$$

$$4. \delta = f_1 + f_2 \left(1 + \frac{r_3 f_2}{r_1 f_1} \right)$$

5. Distance hyperfocale, d'après Centrale TSI 2006

On modélise l'objectif d'un appareil photo par une lentille convergente L , de centre O et de distance focale image f' .

1. Lorsque l'appareil est mis au point à l'infini (objet à l'infini), où doit-on placer la pellicule ?

2. On considère un point objet A à distance d_A devant la lentille. La pellicule est dans la position de la question (1).

a) Exprimer la position $\overline{OA'}$ de l'image A' de A .

b) On note D_L le diamètre utile de la lentille (diamètre du faisceau au niveau de la lentille) et D_A le diamètre de la tache du faisceau sur la pellicule. Montrer que :

$$D_A = D_L \frac{f'}{d_A}$$

c) La pellicule est formée de grains de diamètre ϵ ¹. Une image sera nette tant que l'image du point A aura un diamètre inférieur ou égal à ϵ .

Calculer la distance d_A minimale (appelée distance hyperfocale) qui donnera une image nette sur la pellicule. Application numérique pour $f' = 3,0 \text{ cm}$; $D_L = 2 \text{ mm}$ et $\epsilon = 20 \text{ }\mu\text{m}$

6. Doublet de Huygens

Soient deux lentilles minces $\mathcal{L}_1$ et $\mathcal{L}_2$ de foyers images respectivement F'_1 et F'_2 , de foyers objets respectivement F_1 et F_2 , de distances focales respectives $f'_1 = 3a$ et $f'_2 = a$ et telles que leurs centres respectifs O_1 et O_2 vérifient $\overline{O_1O_2} = e = 2a$. On notera $\Delta = \overline{F'_1F'_2}$.

1. Déterminer graphiquement la position du foyer image F' du système optique formé par $\mathcal{L}_1$ puis $\mathcal{L}_2$.
2. Retrouver ce résultat par un calcul en déterminant l'expression de $\overline{F'_1F'}$.
3. Déterminer graphiquement la position du foyer objet F du système optique formé par $\mathcal{L}_1$ puis $\mathcal{L}_2$.
4. Retrouver ce résultat par un calcul en déterminant l'expression de $\overline{F_1F}$.

7. Lentilles et système afocal, d'après ENAC 2005

1. On dispose un objet $\overline{A_0B_0}$ orthogonalement à l'axe optique d'une lentille divergente $\mathcal{L}_1$ de distance focale $f'_1 = -20 \text{ cm}$. Où doit se trouver l'objet par rapport à la lentille pour que le grandissement transversal soit égal à 0,5?
2. Quelle est alors la position de l'image $\overline{A_1B_1}$?
3. On place après la lentille $\mathcal{L}_1$ un viseur constitué d'une lentille convergente $\mathcal{L}_2$ de même axe optique et de distance focale $f'_2 = 40 \text{ cm}$. On dispose également un écran perpendiculairement à l'axe optique à une distance $\overline{O_2E} = 80 \text{ cm}$ du centre du viseur. Calculer la distance $\overline{O_1O_2}$ entre les deux lentilles pour qu'on obtienne une image sur l'écran de l'objet initial.
4. On souhaite utiliser le dispositif précédent pour transformer un faisceau cylindrique de rayons parallèles à l'axe optique et de diamètre d en un faisceau cylindrique de rayons parallèles à l'axe optique et de diamètre D . Calculer la distance $\overline{O_1O_2}$ entre les deux lentilles pour obtenir un tel résultat.
5. Calculer dans ce cas le rapport entre les deux diamètres $\frac{D}{d}$.

B. Exercices et problèmes

1. Appareil photographique, macrophotographie et téléobjectif

On peut modéliser un appareil photographique par une lentille mince de distance focale $f' = 100 \text{ mm}$. La mise au point s'effectue en déplaçant la lentille par rapport à la pellicule.

¹Dans le cas d'un appareil ou numérique avec un capteur C.C.D., ϵ est la taille d'un pixel (Picture Element)

1. Dans quel intervalle de distances à la pellicule doit-on placer l'objectif pour avoir des images nettes d'objets situés entre 1,4 m et l'infini de cet objectif?
2. Construire l'image $A'B'$ d'un objet AB .
3. Calculer le grandissement pour la photographie d'une tour de 60 m de haut située à 3,0 km ainsi que pour la photographie d'un objet mesurant 1 m à 1,4 m de l'objectif.
4. Macrophotographie.
On place une lentille L_1 de distance focale $f'_1 = 0,333$ m à 5,0 cm devant l'objectif (*id est* la première lentille).

- a) En conservant la course de l'objectif, où sont situés les objets dont ce nouveau dispositif donne des images nettes?
- b) Calculer le grandissement pour un objet de 1,0 m situé à 30 cm de L_1 .
- c) En déduire l'intérêt de la deuxième lentille.

5. Téléobjectif.

On place une lentille L_2 de distance focale $f'_2 = -25$ mm accolée à la première et une lentille L_3 de distance focale $f'_3 = 100$ mm à 75 mm des deux autres.

- a) En conservant la course de l'objectif, où se trouvent les objets dont ce nouveau dispositif donne des images nettes?
- b) Calculer le grandissement de la tour de 60 m située à 3,0 km.
- c) En déduire l'intérêt d'un tel objectif.

2. Doubleur de focale, d'après CCP DEUG 2004

On modélise l'objectif d'un appareil photographique par une lentille convergente dont la distance focale est $f'_1 = 50$ mm. La pellicule se trouve sur un écran plan fixe perpendiculaire à l'axe optique de l'objectif. La mise au point s'obtient en déplaçant l'objectif par rapport à la pellicule sur une distance d variable et ne dépassant pas $d_{\max} = 100$ mm. On s'intéresse à la photographie d'un objet distant de $D = 50$ m de la pellicule et d'une hauteur $h = 2,0$ m.

1. Montrer que la relation de conjugaison permet d'établir une relation entre d , D et f'_1 qu'on mettra sous la forme d'une équation du second degré en d .
2. Calculer alors la valeur prise par d dans le cas envisagé en tenant compte des contraintes pour l'objectif. On fera l'application numérique.
3. Déterminer le grandissement et donner sa valeur numérique.
4. En déduire la taille h' de l'image sur la pellicule et donner sa valeur numérique.
5. On place maintenant entre l'objectif et la pellicule un doubleur de focale assimilable à une lentille divergente de distance focale $f'_2 = -40$ mm à une distance $d_2 = 40$ mm de la pellicule. La distance d' entre l'objectif et la pellicule peut maintenant atteindre $d'_{\max} = 120$ mm. On note AB l'objet à photographier, $A'B'$ son image par l'objectif seul et $A''B''$ l'image finale (formée par l'objectif et le doubleur de focale). $A''B''$ étant sur la pellicule, établir l'expression de d_1 , distance entre $A'B'$ et la pellicule. On fera une détermination par le calcul.
6. Même question par une résolution graphique.

7. Calculer le grandissement γ_2 apporté par le doubleur de focale. Donner sa valeur numérique.
8. L'objet AB étant à une distance D de la pellicule, déterminer la distance d' correspondant à une mise au point correcte. Donner la valeur numérique.
9. Calculer la nouvelle hauteur h'' de l'image sur la pellicule. Donner la valeur numérique.
10. Dédire de tous ces résultats la signification du terme "doubleur de focale".

3. Optique de l'œil, d'après I.C.N.A. 2006

Le cristallin de l'œil est assimilable à une lentille mince de centre optique O , dont la vergence V est variable. L'espace objet est l'air d'indice $n_0 = 1$, et on supposera que l'indice de l'espace image est aussi $n_0 = 1$ (en réalité c'est un milieu assimilable à de l'eau, d'indice $n_1 = 4/3$). L'image se forme sur la rétine, qui dans la réalité est à la distance $d_{\text{rét}} = 15$ mm de O mais que l'on considérera à $d = 11$ mm en raison de la différence entre n_0 et n_1 .

1. Un observateur doté d'une vision "normale" regarde un objet AB placé à 1 m devant lui, et tel que $\overline{AB} = 10$ cm.
 - a) Préciser si l'image est réelle ou virtuelle, droite ou renversée.
 - b) On note $A'B'$ l'image de AB sur la rétine. Calculer le grandissement $\gamma = \frac{\overline{A'B'}}{\overline{AB}}$, et en déduire la taille de l'image $\overline{A'B'}$.
 - c) Calculer la vergence V du système.
2. L'observateur regarde maintenant un objet placé à 25 cm devant lui.
 - a) Préciser si l'image est réelle ou virtuelle, droite ou renversée.
 - b) Calculer la variation de la vergence par rapport à celle de la question (c) ainsi que la taille de l'image.
3. On s'intéresse maintenant à un individu myope, qui possède donc un cristallin trop convergent. Lorsqu'il regarde à l'infini, l'image se forme à 0,5 mm en avant de la rétine (située à $d = 11$ mm de O). Pour corriger ce problème, cette personne est dotée de lunettes dont chaque verre est assimilé à une lentille mince de vergence V' constante et de centre optique O' , placé à $\ell = 2$ cm de O .
 - a) Calculer la vergence V' des verres de lunettes.
 - b) L'individu observe un objet situé à 1 m devant lui. Calculer la position de l'image intermédiaire ainsi que le grandissement de l'ensemble (lunette-cristallin).

4. Étude d'une lunette astronomique, d'après ENSTIM 2004

Une lunette astronomique est schématisée par deux lentilles minces convergentes de même axe optique Δ :

- l'une L_1 (objectif) de distance focale image $f'_1 = \overline{O_1F'_1}$;
- l'autre L_2 (oculaire) de distance focale image $f'_2 = \overline{O_2F'_2}$.

On rappelle qu'un œil normal voit un objet sans accommoder si celui-ci est placé à l'infini.

On souhaite observer la planète Mars qui est vue à l'œil nu sous un diamètre apparent α .

1. Pour observer la planète avec la lunette, on forme un système afocal.

a) Que cela signifie-t-il ? En déduire la position relative des deux lentilles.

b) Faire le schéma de la lunette pour $f'_1 = 5f'_2$.

Dessiner sur ce schéma la marche à travers la lunette d'un faisceau lumineux (non parallèle à l'axe) formé de rayons issus de l'astre. On appelle $\overline{A'B'}$ l'image intermédiaire.

c) On souhaite photographier cette planète. Où faut-il placer la pellicule ?

2. On note α' , l'angle que forment les rayons émergents extrêmes en sortie de la lunette.

a) L'image est-elle droite ou renversée ?

b) La lunette est caractérisée par son grossissement $G = \frac{\alpha'}{\alpha}$. Exprimer G en fonction de f'_1 et f'_2 .

c) Le principal défaut d'une lentille est appelé défaut d'aberrations chromatiques : expliquer brièvement l'origine de ce défaut et ses conséquences. Pour quelle raison un miroir n'a-t-il pas ce défaut ?

3. On veut augmenter le grossissement de cette lunette et redresser l'image. Pour cela, on interpose entre L_1 et L_2 , une lentille convergente L_3 de distance focale image $f'_3 = \overline{O_3F'_3}$. L'oculaire L_2 est déplacé pour avoir de la planète une image nette à l'infini à travers le nouvel ensemble optique.

a) Quel couple de points doit conjuguer L_3 pour qu'il en soit ainsi ?

b) On appelle γ_3 , le grandissement de la lentille L_3 . En déduire $\overline{O_3F'_1}$ en fonction de f'_3 et γ_3 .

c) Faire un schéma (On placera O_3 entre F'_1 et F_2 et on appellera $\overline{A'B'}$ la première image intermédiaire et $\overline{A''B''}$, la seconde image intermédiaire).

d) En déduire le nouveau grossissement G' en fonction de G et γ_3 . Comparer G' à G en signe et valeur absolue.

5. La loupe, d'après ENM Saint-Cyr 1994

Dans tout le problème on s'intéresse à des systèmes centrés.

Un observateur emmétrope, c'est à dire ayant un œil normal peut voir distinctement de l'infini à une distance minimale d_m . On dit que l'observateur accommode si l'objet qu'il observe n'est pas à l'infini.

Cet observateur regarde à l'œil nu un tout petit objet plan que l'on assimilera à un segment AB de longueur ℓ , perpendiculaire à l'axe optique.

1. Déterminer α_m , l'angle maximal sous lequel l'objet peut-être vu.
2. L'observateur regarde AB à travers une lentille mince convergente de distance focale f' et de centre O (loupe). Son œil est situé à une distance a de la loupe ($a < d_m$).
- a) Déterminer les positions de l'objet rendant possible l'observation d'une image nette par l'observateur emmétrope. Faire une construction géométrique de l'image. L'image est-elle droite ou renversée ?
- b) Pour quelle position de l'objet, l'observation se fait-elle sans accommodation ? Exprimer l'angle α sous lequel l'œil voit l'image. Application numérique : que vaut le grossissement commercial de la loupe $G = \alpha/\alpha_m$? On donne $d_m = 0,25$ m et $f' = 50$ mm.

6. Principe du viseur, d'après ENM Saint-Cyr 1994

Ce problème est la suite du problème 5. qui doit être traitée auparavant.

Un viseur est constitué d'un objectif et d'un oculaire de même axe optique (Ox). On assimilera l'objectif à une lentille mince convergente L_1 de centre O_1 et de distance focale f'_1 et l'oculaire à une lentille mince convergente L_2 de centre O_2 et de distance focale f'_2 . On pose $\overline{O_1O_2} = D$ et $\overline{OO_2} = d$ (les distances D et d sont positives et réglables). Dans un plan transverse, est disposé en O une croix appelée réticule, constituée de deux traits fins perpendiculaires. L'observateur place son œil à une distance a derrière l'oculaire ($a < d_m$).

1. Quel est l'intervalle des valeurs de d permettant à l'observateur de voir le réticule net à travers l'oculaire ? Où doit-on placer le réticule pour l'observer sans accommodation ? On effectuera les constructions dans les deux cas extrêmes de d .
2. Le réglage précédent est supposé réalisé. On souhaite observer à travers le viseur, un objet A situé sur l'axe optique à l'abscisse $x = \overline{OA}$. L'image intermédiaire doit se trouver dans le plan du réticule.
- a) Montrer que $x \leq -4f'_1$
- b) Déterminer l'expression de D en fonction de x .
- c) En déduire la plage de réglage de la distance D que le constructeur doit prévoir.
3. a) Un observateur myope souhaite utiliser le viseur sans ses verres correcteurs pour observer un objet A situé à l'infini, dans les conditions définies précédemment. Sachant que sa distance maximale de vision distincte est δ , calculer les valeurs des réglages qu'il doit effectuer.
- b) En supposant que tous les utilisateurs du viseur, qu'ils soient myopes ou hypermétropes, ont des verres correcteurs de vergence comprise entre -8δ et 8δ , déterminer la position de l'objet que l'œil peut voir sans accommoder et sans verres (on supposera lors de leur utilisation que les verres correcteurs sont accolés à l'œil).

REMARQUE : une fois corrigés, les yeux sont supposés emmétropes.

Déterminer la plage de réglage de l'oculaire à prévoir pour que le viseur soit utilisable par tous sans verres correcteurs. Données : $a = 0$ et $f'_2 = 2$ cm.

12

Instruments de Travaux Pratiques

1. L'œil

Avant d'étudier les instruments utilisés en travaux pratiques, il est nécessaire d'avoir quelques connaissances sur le fonctionnement et la modélisation de l'œil. Il n'est pas question ici d'aborder des notions de biologie mais plutôt d'étudier l'optique de l'œil.

1.1 Description

Une coupe de l'œil est présentée sur la figure 12.1.

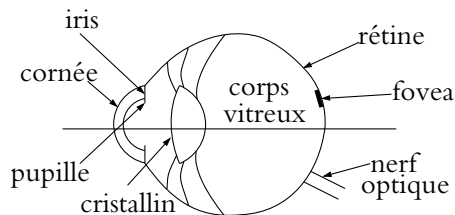


Figure 12.1 Structure de l'œil.

- L'iris (partie colorée) est percée de la pupille dont le diamètre est variable (de 2 à 8 mm) et qui joue le rôle de diaphragme c'est-à-dire qui limite l'intensité lumineuse pénétrant dans l'œil.
- Le cristallin est un muscle assimilable à une lentille mince biconvexe dont la distance focale est variable selon sa contraction. Il donne d'un objet une image renversée sur la rétine.
- La rétine est constituée de cellules sensibles à la lumière (les cônes et les bâtonnets).

- La fovéa est une petite dépression non située sur l'axe optique et entourée de la macula (tache jaune) sur laquelle on observe au mieux les petits détails de l'image. D'ailleurs, les astronomes amateurs savent que pour observer des objets de faible intensité à l'œil nu il ne faut pas les regarder de face mais légèrement de côté.
- Le corps vitreux est formé d'une substance gélatineuse d'indice de réfraction $n = 1,336$.
- L'ensemble rétine-nerf optique code l'image sous forme d'influx nerveux et l'envoie au cerveau par l'intermédiaire du nerf optique. Le cerveau interprète le message (retournement de l'image, correction de la distorsion, vision stéréoscopique, c'est-à-dire impression de relief par combinaison des informations provenant des deux yeux). En ce qui concerne la vision stéréoscopique, on peut douter de l'utilité des deux yeux car lorsqu'on se cache un œil, on a toujours l'impression de voir les objets en relief. Cela vient de la mémoire que l'on a des objets. Si on ne possédait qu'un œil, on ne pourrait voir en relief.

L'ensemble de l'œil sera modélisé par une lentille mince convergente formant une image sur la rétine modélisée par un écran. Il faut être conscient que ceci est une approximation.

1.2 Caractéristiques

On imagine un cône partant de la pupille tel que tout objet se trouvant dans ce cône pourra avoir une image nette sur la rétine. On appelle *champ angulaire* l'angle de ce cône de vision. Le champ angulaire a une valeur importante (40 à 50°), cependant la zone de perception des détails est beaucoup plus réduite (1°), car l'image doit alors se former sur la fovéa.

L'œil ne distingue deux détails différents de l'objet que si **leur image se forme sur deux cellules différentes de la rétine** (Cf. le grain d'une pellicule photo). Dans de bonnes conditions d'éclairement (ni trop sombre, ni trop lumineux), l'œil distingue des détails d'environ $1'$ d'arc, soit 3.10^{-4} rad. Cette valeur constitue la *limite de résolution* ou *pouvoir séparateur* de l'œil. La résolution pratique dépend fortement des conditions d'éclairement et du contraste si bien que la limite précédente est rarement atteinte.

L'œil ne peut voir une image nette que si elle se forme sur la rétine. Un œil au repos (cristallin non contracté) voit à une distance maximale D_m . Le point situé à cette distance porte le nom de *Punctum Remotum* noté P.R. Pour voir des objets plus proches, le cristallin doit se contracter pour être plus convergent (sa vergence augmente), on dit que l'œil *accommode*. Le point le plus proche que peut voir net un œil est appelé *Punctum Proximum* noté P.P. Il correspond à la distance minimale de mise au point. La zone située entre le P.P. et le P.R. est appelé *champ en profondeur de l'œil*.

Pour un œil normal, le P.R. est à l'infini et le P.P. à environ 25 cm pour un adulte. Le P.P. est plus proche pour un enfant, c'est pour cela que souvent les enfants lisent avec « le nez sur leur livre ».

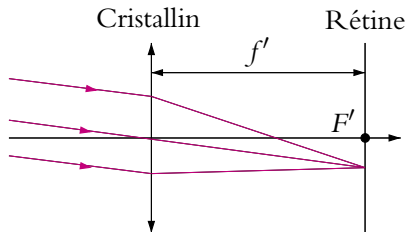


Figure 12.2 Formation de l'image d'un objet au P.R.

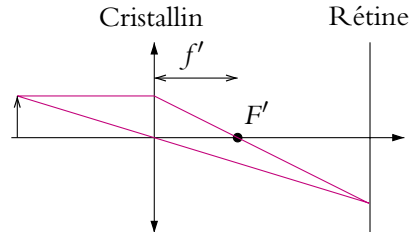


Figure 12.3 Formation de l'image d'un objet au P.P.

Les figures 12.2 et 12.3 représentent respectivement la formation de l'image d'un objet au P.R. et au P.P.

1.3 Les défauts de l'œil

Un œil normal est dit *emmétrope*. Voici quelques défauts principaux des yeux couramment rencontrés.

a) La myopie

Un œil *myope* possède un cristallin trop convergent. Par conséquent, le P.P. est plus proche que pour l'œil normal et le P.R. n'est plus à l'infini. On s'aperçoit sur la figure 12.4 que l'image d'un objet ponctuel à l'infini se forme avant la rétine et l'observateur ne voit qu'une tache floue (le myope sans ses lunettes voit flou de loin). Il faut rendre l'œil moins convergent donc la correction se fait par l'utilisation d'une lentille divergente.

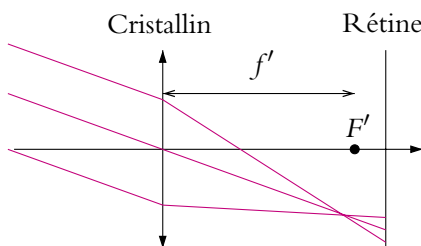


Figure 12.4 Vision d'un objet à l'infini pour un œil myope.

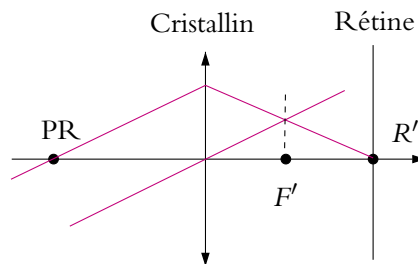


Figure 12.5 Position du P.R. pour un œil myope.

Pourquoi voit-on les yeux d'un myope plus petits derrière ses lunettes que sans ses lunettes ? L'œil du myope joue alors le rôle d'objet réel pour la lentille de la lunette. Si l'on se réfère à la construction de la figure 11.15 du paragraphe sur les lentilles, on s'aperçoit que dans ce cas, l'image est virtuelle et plus petite, ce que l'on observe.

La figure 12.5 montre comment on détermine la position du P.R. d'un myope. On cherche l'antécédent d'une image R' se formant sur la rétine dans le cas de l'œil au repos. Pour cela on cherche le rayon incident croisant l'axe dont le rayon émergent arrive en R' . On trace le rayon parallèle passant par le centre optique, les deux rayons se croisent dans le plan focal image.

b) L'hypermétropie

Un œil *hypermétrope* possède un cristallin pas assez convergent. L'œil doit donc accommoder pour voir à l'infini sinon l'image se forme derrière la rétine (Cf. figure 12.6) et l'observateur observe une tache floue. Le P.P. est plus éloigné que pour l'œil normal. Il faut rendre l'œil plus convergent donc la correction se fait par l'utilisation d'une lentille convergente.

Au contraire des myopes, pourquoi voit-on les yeux des hypermétropes plus grands derrière leurs lunettes que sans leurs lunettes ? Le raisonnement est semblable au précédent mais il faut se référer à la construction de la figure 11.13 du paragraphe sur les lentilles. C'est le cas de la loupe : l'image est virtuelle et plus grande, ce que l'on observe.

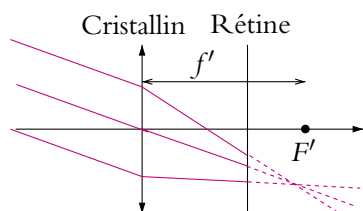


Figure 12.6 Vision d'un objet à l'infini pour un œil hypermétrope.

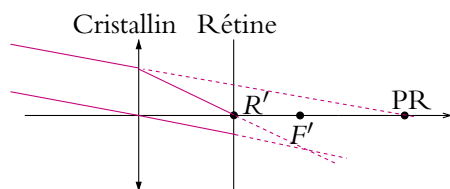


Figure 12.7 Position du P.R. d'un œil hypermétrope.

Comme pour le myope, on peut déterminer la position du P.R. de l'hypermétrope (Cf. figure 12.7). Il se trouve derrière l'œil, ce qui fait dire qu'un hypermétrope « peut voir derrière lui sans se retourner ». C'est bien sûr faux, la construction montre seulement qu'un faisceau qui convergerait au P.R. donne une image nette sur la rétine.

c) L'astigmatisme

L'œil *astigmatique* ne possède pas la symétrie de révolution. La correction se fait par des lentilles non sphériques. On peut simuler un œil astigmatique en inclinant une lentille par rapport à l'axe optique. Cette expérience a été décrite dans le paragraphe sur les aberrations.

d) La presbytie

L'œil *presbyte* a des difficultés à accommoder. C'est un problème qui est dû au vieillissement, le cristallin perdant progressivement sa plasticité. Les objets proches sont

moins bien vus et surtout le champ en profondeur diminue. La correction se fait soit avec des verres multifocaux (une correction pour la vision de loin et une correction pour la vision de près), soit avec des verres progressifs (c'est-à-dire de focale variable selon la distance œil-objet).

2. Sources de lumière

Avant d'étudier plus en détail les instruments de travaux pratiques, il faut avoir quelques informations sur les sources de lumières qu'on utilise. C'est le but de ce paragraphe.

2.1 Dispositif expérimental de visualisation des spectres par projection

Une lumière contient en général des radiations de longueurs d'onde différentes. On peut décomposer la lumière en la faisant passer à travers un système dispersif, par exemple un prisme. En effet, en raison de la loi de Cauchy citée dans le paragraphe sur les aberrations des lentilles, le prisme dévie la lumière plus ou moins en fonction de la longueur d'onde : **le bleu est plus dévié que le rouge**.

La figure obtenue est appelée *spectre*.

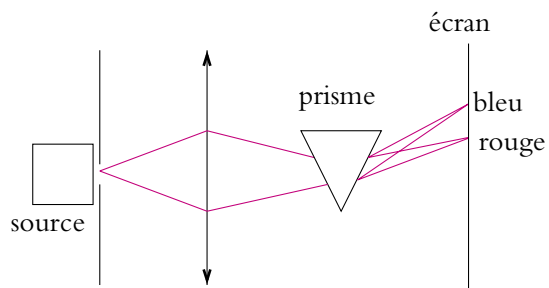


Figure 12.8 Dispersion de la lumière blanche par un prisme.

Dans le montage proposé sur la figure 12.8, on réalise l'image d'une fente sur l'écran, à l'aide d'une lentille convergente. En l'absence de prisme, on obtiendrait sur l'écran une image unique de la fente, de la couleur de la source. Le prisme décompose la lumière et donne une image de la fente pour chaque couleur présente dans la lumière émise par la source.

2.2 Source de lumière blanche

Il s'agit des *lampes à incandescence*. Un filament de tungstène chauffé émet une onde électromagnétique dont une partie importante se trouve dans le spectre visible. On parle de *rayonnement thermique*. Le spectre d'émission est continu (à opposer aux spectres de raies des sources présentées ultérieurement). En fait, la lumière n'est pas parfaitement blanche (bien qu'elle le paraisse) c'est-à-dire que toutes les longueurs d'ondes ne sont pas émises avec la même intensité. La longueur d'onde $\lambda_{\max}$ du maximum d'émission peut être obtenue approximativement par la loi de Wien (Cf. programme de deuxième année MP) :

$$\lambda_{\max} T = 2,987 \cdot 10^{-3} \text{ K.m} = 2987 \text{ }\mu\text{m.K} \quad (12.1)$$

où T est la température du corps exprimée en kelvins.

On distingue deux types de lampes :

- les lampes classiques,
- les lampes halogènes.

Dans le cas des lampes classiques, la température du filament est de 2 900 K. L'atmosphère dans laquelle baigne le filament est neutre (par exemple du krypton). D'après la relation (12.1), le maximum d'émission est dans l'infrarouge. Beaucoup d'énergie est perdue sous forme de chaleur (rayonnée dans l'infrarouge) et la lumière apparaît plutôt jaune. Le rendement d'une telle lampe est donc faible.

Pour translater le maximum d'émission de l'infrarouge au visible et le rapprocher du maximum d'émission solaire ($0,56 \mu\text{m}$), il faut, d'après la loi de Wien (12.1), augmenter la température du filament. Mais dans ce cas, le tungstène se sublime (passage de l'état solide à l'état gazeux), de nombreux atomes de tungstène sont perdus par le filament et ont tendance à se déposer sur la paroi de l'ampoule et le filament finit par casser rapidement.

Aussi, dans le cas des lampes halogènes, l'atmosphère dans laquelle baigne le filament n'est pas inerte (en général de l'iode). Le tungstène vaporisé se combine alors avec l'iode puis se redépote sur le filament et le régénère. La température pouvant être atteinte par l'ampoule est supérieure à la précédente, l'enveloppe en verre doit alors être remplacée par du quartz (plus résistant à la chaleur) : on parle d'*ampoule quartz-iode*.

2.3 Lampes spectrales (tube à décharge)

Les électrons d'un atome ne peuvent se trouver dans un état d'énergie quelconque. Les niveaux d'énergie sont quantifiés (Cf. figure 12.9).

À l'équilibre, les électrons se trouvent dans une configuration telle que l'énergie de l'atome est minimale, appelée *état fondamental*. Si on communique de l'énergie à un gaz, par exemple sous forme électrique dans un tube à décharge, les atomes du gaz passent à un état d'énergie supérieure, appelé *état excité*. Cet état n'est cependant pas stable, et l'atome revient spontanément à l'état fondamental en rendant l'énergie excédentaire sous forme de lumière (émission d'un photon). C'est ce qu'on appelle l'*émission spontanée*. Si on utilise les notations suivantes :

- E_0 : énergie du niveau fondamental,
- E_1 : énergie du niveau excité,
- h : constante de Planck ($h = 6,62610^{-34} \text{ J.s}$),

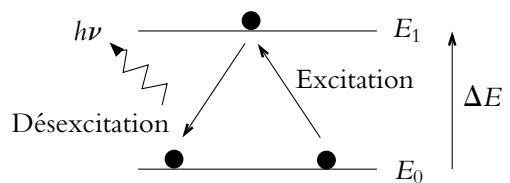


Figure 12.9 Principe de l'émission spontanée.

- ν : fréquence du photon émis,

on a la relation suivante :

$$\Delta E = E_1 - E_0 = h\nu \quad (12.2)$$

On obtient un spectre de raies (non continu), caractéristique du gaz.

Par exemple :

- Hydrogène : série de Balmer (voir cours de chimie) ;
- Sodium : principalement doublet jaune à 589,0 et 589,6 nm ;
- Mercure : raies les plus visibles :
 - rouge : 623,4 nm
 - doublet jaune : 579,1 et 577,0 nm
 - vert : 546,1 nm
 - bleu : 435,8 nm
 - violet : 404,7 nm
- Cadmium : raies les plus visibles :
 - rouge : 643,8 nm
 - vert : 508,6 nm
 - bleu : 480,0 nm

2.4 Laser

Il s'agit de l'abréviation de *Light Amplifier by Stimulated Emission of Radiation* (amplificateur de lumière par émission stimulée). On a parlé précédemment d'émission spontanée ; il est possible aussi de provoquer une émission par une onde incidente rencontrant les atomes : on parle d'*émission stimulée*. Dans ce cas, la longueur d'onde émise est celle de l'onde incidente. On obtient une **lumière quasi monochromatique**, ce qui n'est pas le cas de l'émission spontanée. D'autre part, on utilise des systèmes optiques tels que le faisceau obtenu soit très peu divergent.

Le principe du laser à gaz est le suivant :

On fabrique une cavité afocale (Cf. figure 12.10) à l'aide de deux miroirs sphériques, dont les foyers sont confondus, et entourant un tube à décharge. Grâce à ce système, un rayon revient sur lui-même après quatre réflexions et stimule l'émission du gaz. L'un des miroirs est légèrement transparent pour qu'une partie de la lumière puisse sortir de la cavité et être utilisée.

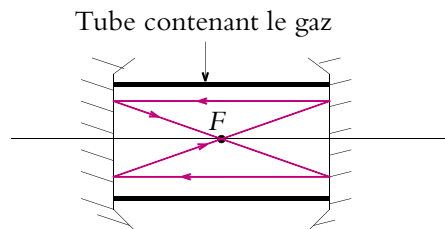


Figure 12.10 Cavité laser.

Les lasers à gaz les plus courants sont les lasers Hélium-Néon dont la longueur d'onde est 632,8 nm soit une émission dans le rouge mais il en existe des verts, jaunes et oranges. On rencontre aussi d'autres types de lasers : à colorant (toutes couleurs), diodes laser (solide), etc., utilisés par exemple dans les lecteurs de disques compacts.

3. Loupes et oculaires

On raisonne dans toute la suite avec l'œil emmétrope.

3.1 La loupe

Plus on approche un objet du P.P. de l'œil, plus ce dernier se fatigue rapidement. Or, pour pouvoir observer les détails d'un objet, il est nécessaire de placer cet objet au P.P. D'autre part, il a été vu que le pouvoir de résolution angulaire de l'œil moyen est au maximum de 3.10^{-4} rad, ce qui permet d'observer au mieux des détails d'environ 1/10 de mm à 30 cm dans des conditions idéales d'éclairement.

Pour observer des détails plus petits ou pour éviter de fatiguer l'œil rapidement, il est nécessaire d'utiliser un instrument grossissant. Le plus simple est de prendre une lentille convergente et de placer l'objet entre le foyer objet et la lentille. Dans ce cas, comme le montre la figure 12.11, on obtient une image virtuelle droite plus grande que l'objet. Si l'on place l'œil au foyer image de la lentille, on se rend compte sur la figure 12.11 que si on déplace l'objet l'image sera toujours observée sous le même angle α .

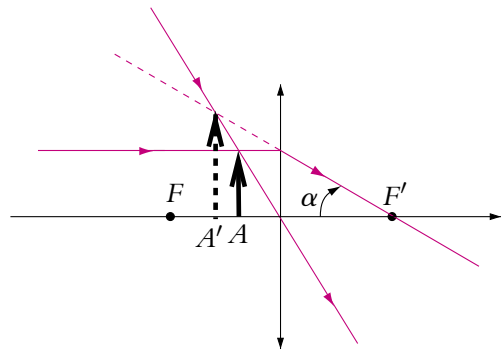


Figure 12.11 Image par une loupe.

3.2 L'oculaire

Chacun a pu cependant se rendre compte, en utilisant une loupe, que le champ d'observation est assez limité et que l'image est fortement distordue sur les bords. Ainsi, pour les instruments d'optique plus sophistiqués, on a recours à un oculaire qui est formé de plusieurs lentilles (souvent deux). Parmi les oculaires à deux lentilles (L_1 de centre O_1 et de distance focale image f'_1 et L_2 de centre O_2 et de distance focale image f'_2) les plus couramment utilisés, on peut citer :

- l'oculaire de Ramsden avec $\frac{f'_1}{3} = \frac{\overline{O_1 O_2}}{2} = \frac{f'_2}{3}$;
- l'oculaire d'Huygens avec $\frac{f'_1}{3} = \frac{\overline{O_1 O_2}}{2} = f'_2$.

L'oculaire d'Huygens a un foyer objet virtuel situé entre les deux lentilles alors que celui de Ramsden est réel (avant la première lentille).

L'oculaire permet d'obtenir une qualité d'image meilleure qu'avec une seule lentille. L'association de plusieurs lentilles permet d'atténuer les aberrations chromatiques, d'obtenir une distorsion plus faible et un champ plus grand.

Comme son nom l'indique, dans un instrument d'optique, **l'oculaire est la partie de l'instrument qui est du côté de l'œil** tandis que **l'objectif est du côté de l'objet que l'on veut observer**. Ainsi l'oculaire est le dernier système optique que rencontrent les rayons avant de pénétrer dans l'œil. On étudiera plus en détails dans un paragraphe ultérieur l'instrument en entier mais il est utile de préciser dès maintenant que l'objectif donne de l'objet observé une image intermédiaire qui, elle, sert d'objet pour l'oculaire.

La première lentille de l'oculaire (du côté de l'objectif) est appelée *lentille ou verre de champ* alors que la dernière lentille (du côté de l'œil) est appelée *lentille ou verre d'œil*. On nomme F_{oi} le foyer objet de la lentille d'œil et F_{oc} le foyer objet de l'oculaire complet. Un objet AB placé dans le plan focal objet de l'oculaire aura une image finale à l'infini de même qu'un objet placé dans le plan focal de la lentille d'œil. D'autre part, comme le montre le schéma de la figure 12.12 réalisé avec deux lentilles, l'image intermédiaire $A'B'$ de AB se trouve dans le plan focal objet de la lentille d'œil puisque l'image finale est à l'infini :

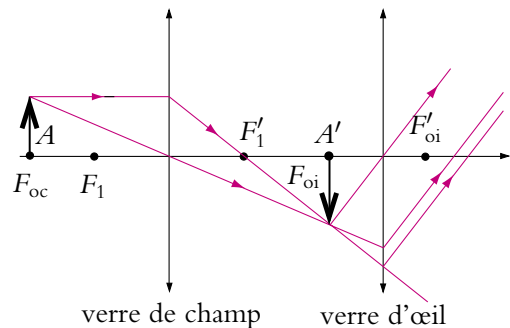


Figure 12.12 Image par un oculaire.

$$AB \xrightarrow{\text{verre champ}} A'B' \xrightarrow{\text{verre œil}} \infty$$

Ce résultat est généralisable au cas de plus de deux lentilles.

Un objet placé dans le plan focal objet de l'oculaire et un objet placé dans le plan focal objet de la lentille d'œil auront leurs images superposées à l'infini.

Comment régler l'oculaire pour éviter une fatigue de l'œil ? Pour éviter toute accommodation, il est nécessaire que l'image finale donnée par l'oculaire soit au P.R. de l'œil donc à l'infini pour un œil emmétrope. L'objet observé par l'oculaire (image intermédiaire de l'instrument) est alors au foyer objet de l'oculaire. On dit dans ce cas que l'oculaire est *régulé à l'infini*.

Si on veut superposer une règle à l'image pour effectuer des mesures, on place cette règle soit dans le plan focal de l'oculaire, soit d'après ce qui a été vu précédemment dans le plan focal du verre d'œil.

4. Lunettes et viseurs

Comme cela a été précisé précédemment, un instrument réglé correctement évite toute accommodation de l'œil. L'image finale donnée par l'instrument doit être au P.R. de l'œil donc à l'infini pour un œil emmétrlope.

Aussi dans toute la suite, on considérera que l'image finale est à l'infini.

Il existe trois types de mesures en travaux pratiques :

- les mesures d'angles pour lesquelles l'objet est à l'infini : on utilise une *lunette afocale* appelée aussi *lunette de visée à l'infini* donnant d'un objet à l'infini une image à l'infini, c'est le principe d'une lunette astronomique ;
- les mesures de distances longitudinales c'est-à-dire le long de l'axe optique : on utilise une lunette appelée *viseur* ou *lunette à frontale fixe* qui, d'un objet à distance finie donne une image à l'infini ;
- les mesures de distances transversales (mesures de grandissement) c'est-à-dire perpendiculairement à l'axe optique : on utilise soit un viseur avec une règle micrométrique ou millimétrique placée dans l'oculaire comme expliqué au paragraphe précédent, soit un viseur monté sur une crémaillère perpendiculaire à l'axe optique.

Dans la suite, l'objectif sera modélisé par une lentille unique.

4.1 Lunette de visée à l'infini ou afocale

a) Schéma de principe

L'objet à l'infini a une image intermédiaire $A_i B_i$ dans le plan focal image de l'objectif. Pour que l'image de $A_i B_i$ par l'oculaire soit à l'infini, il faut que $A_i B_i$ soit dans le plan focal objet de l'oculaire (Cf. figure 12.13). **La lunette est donc réglée à l'infini si le foyer objet de l'oculaire est confondu avec le foyer image de l'objectif :**

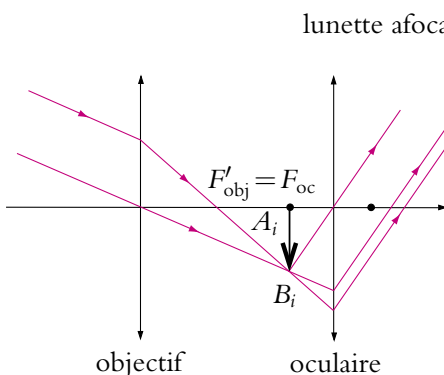


Figure 12.13 Image par une lunette.

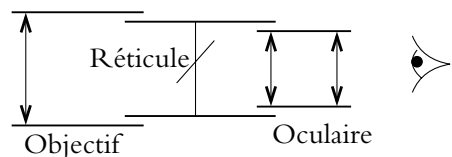


Figure 12.14 Structure d'une lunette.

On peut remarquer sur la figure 12.13 que, sans la lunette, l'œil aura l'impression que les rayons proviennent du haut et qu'avec la lunette, il verra les rayons provenir du bas : **l'image est donc renversée.**

Une lunette est formée de trois tubes (Cf. figure 12.14) : le premier contenant l'objectif, le dernier l'oculaire (représenté ici avec ses deux lentilles). Ces deux tubes sont mobiles par rapport au tube du milieu contenant un repère en forme de croix que l'on nomme *réticule*, ici représenté en perspective.

4.2 Réglage d'une lunette afocale

Le réglage de la lunette a pour but de faire coïncider le plan focal objet de l'oculaire et le plan focal image de l'objectif donc de réaliser $F_{oc} = F'_{obj}$ où F_{oc} est le foyer objet de l'oculaire et F'_{obj} le foyer image de l'objectif. Pour cela, il faut disposer d'un repère qui n'est autre que le réticule. Le réglage s'effectue en deux étapes :

	But du réglage	Mode opératoire
Réglage de l'oculaire.	Faire correspondre le plan du réticule et le plan focal objet de l'oculaire, de manière à ce que l'image du réticule par l'oculaire soit à l'infini.	Tourner la bague de l'oculaire jusqu'à voir le réticule net. On vérifie que le réglage est correct en approchant rapidement l'œil de l'oculaire. Celui-ci doit voir immédiatement le réticule net. S'il y a un léger temps d'accommodation, le réglage est à refaire.
Réglage de l'objectif.	Faire correspondre le plan du réticule et le plan focal image de l'objectif, de manière à ce que l'image de l'objet par l'objectif soit dans le plan du réticule (qui est aussi plan focal objet de l'oculaire).	Viser un objet à grande distance. Tourner la bague de l'objectif jusqu'à voir l'image et le réticule nets simultanément. Pour vérifier le réglage, effectuer un léger hochement de tête devant l'oculaire. Si l'image observée reste immobile par rapport au réticule, le réglage est correct. Si l'image semble se déplacer par rapport au réticule, parfaire le réglage.

Le réglage est alors terminé et la lunette se trouve dans la configuration représentée sur la figure 12.13.

On peut se poser trois questions :

- Pourquoi ne pas effectuer le réglage une fois pour toute au moment de la fabrication de l'instrument et rendre solidaires les trois tubes ?
- Faut-il modifier le réglage quand l'utilisateur change ?
- Pourquoi faut-il hocher la tête pour vérifier le réglage ?

La deuxième question répond en partie à la première. Si la lunette était réglée à l'infini une fois pour toute, son réglage ne conviendrait qu'à un utilisateur emmétrope. De plus, les manipulations risqueraient au fil du temps de modifier les positions relatives des tubes et de dérégler la lunette.

Quand l'utilisateur change, **il ne faut absolument pas toucher au réglage de l'objectif mais on modifie le réglage de l'oculaire**. Il suffit de conjuguer le plan du réticule (qui est celui de l'image intermédiaire) avec le P.R. de l'œil donc de tourner la bague de l'oculaire jusqu'à voir le réticule net sans accommodation.

En ce qui concerne le hochement de tête, c'est une méthode qui permet de corriger l'erreur dite de *parallaxe*. On veut que l'image intermédiaire soit exactement dans le plan du réticule. Mais l'œil peut voir simultanément nets des objets qui sont dans des plans proches mais différents : c'est ce qu'on appelle *la latitude de mise au point*. Cependant si on bouge la tête, on se rendra compte que les objets ne sont pas dans le même plan s'ils se déplacent l'un par rapport à l'autre, ce qui n'est pas le cas s'ils sont dans le même plan. Le lecteur peut s'en convaincre en regardant autour de lui deux objets dans des plans différents et en remuant la tête.

4.3 Lunettes de visée autocollimatrices

Dans le mode opératoire proposé précédemment, il est nécessaire de pouvoir viser un objet à grande distance pour effectuer le réglage de la lunette. Ce n'est pas toujours possible quand on se trouve dans un laboratoire. Pour pallier ce problème, une lunette dite *autocollimatrice* existe avec la possibilité de fabriquer un objet à l'infini.

En plus d'une lunette classique, cette lunette dispose d'une ampoule et d'une lame semi-réfléchissante dans le tube intermédiaire (Cf. figure 12.15). Une telle lame réfléchit une fraction de la lumière mais laisse passer l'autre fraction. Si la lame est descendue comme sur la figure, l'ampoule éclaire le réticule, qui peut cependant toujours être vu à travers la lame puisqu'elle est semi-réfléchissante.

On a dit précédemment que l'objectif est réglé lorsque le réticule se trouve dans son plan focal image. On utilise un mode opératoire permettant de déterminer rapidement le plan focal d'un système convergent : c'est *l'autocollimation* qui est décrite dans le chapitre sur la focométrie au paragraphe 3.

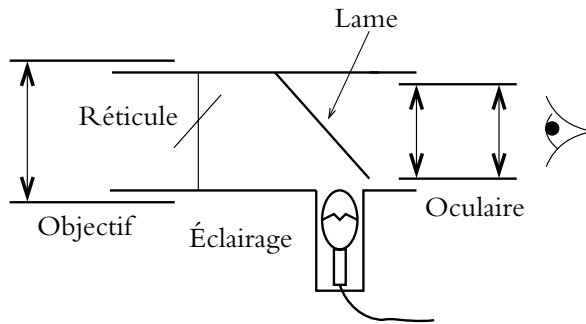


Figure 12.15 Structure d'une lunette autocollimatrice.

Une fois l'oculaire réglé, on éclaire le réticule grâce à la lame semi-réfléchissante et on place un miroir plan devant l'objectif. Le réticule $\mathcal{R}$ éclairé sert d'objet qui va donner une image $\mathcal{R}'$ à travers le système objectif-miroir. On appelle $\mathcal{R}_1$ et $\mathcal{R}_2$ les images intermédiaires :

$$\mathcal{R} \xrightarrow{\text{objectif}} \mathcal{R}_1 \xrightarrow{\text{miroir}} \mathcal{R}_2 \xrightarrow{\text{objectif}} \mathcal{R}$$

Attention, puisque la lumière change de sens entre $\mathcal{R}_2$ et $\mathcal{R}$, il faut penser dans les raisonnements à inverser les foyers objet et image sachant que, puisque l'objectif est assimilable à une lentille mince, les deux foyers sont symétriques par rapport à la lentille.

On suppose que $\mathcal{R}$ est dans le plan focal objet de l'objectif alors $\mathcal{R}_1$ sera à l'infini de même que $\mathcal{R}_2$ et $\mathcal{R}'$ sera alors dans le plan focal devenu le plan focal image. Ainsi $\mathcal{R}$ et $\mathcal{R}'$ seront dans le même plan et on les verra **simultanément nets**.

On peut donc récapituler le réglage d'une lunette autocollimatrice.

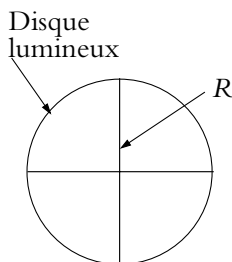


Figure 12.16 Observation de l'oculaire réglé.

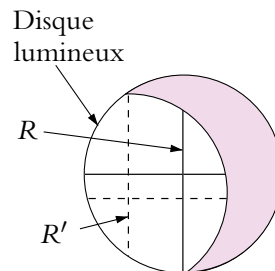


Figure 12.17 Observation de la lunette réglée.

	But du réglage	Mode opératoire
Réglage de l'oculaire. Cf. figure 12.16	Faire correspondre le plan du réticule et le plan focal objet de l'oculaire, de manière à ce que l'image du réticule par l'oculaire soit à l'infini.	Tourner la bague de l'oculaire jusqu'à voir le réticule net. On vérifie que le réglage est correct en approchant rapidement l'œil de l'oculaire. Celui-ci doit voir immédiatement le réticule net. S'il y a un léger temps d'accommodation, le réglage est à refaire.
Réglage de l'objectif par autocollimation. Cf. figure 12.17	Faire correspondre le plan du réticule et le plan focal image de l'objectif, de manière à ce que l'image de l'objet par l'objectif soit dans le plan du réticule (qui est aussi plan focal objet de l'oculaire).	Abaisser la lame semi-réfléchissante et éclairer le réticule. Placer un miroir plan devant l'objectif. Tourner la bague de l'objectif jusqu'à voir le réticule et son image nets simultanément. Pour vérifier le réglage, effectuer un léger hochement de tête devant l'oculaire. Si l'image observée reste immobile par rapport au réticule, le réglage est correct. Si l'image semble se déplacer par rapport au réticule, parfaire le réglage.

► **Remarque :** Pour la plupart des lunettes, la lame semi-réfléchissante est escamotable. Une fois la lunette réglée, on peut relever cette lame pour éviter d'avoir des réflexions parasites lors des observations.

Lorsque l'utilisateur est myope ou hypermétrope, il suffit de modifier le réglage de l'oculaire. Pour un myope, le P.R. étant devant l'œil, il faut que l'oculaire donne une image virtuelle comme une loupe : il faut donc rapprocher l'oculaire du réticule. Par contre, pour un hypermétrope, le P.R. étant derrière l'œil, il faut que l'oculaire donne une image réelle : il faut donc éloigner l'oculaire du réticule.

4.4 Viseurs ou lunettes à frontale fixe

Le viseur est une lunette donnant d'un objet à distance finie une image à l'infini. Comment peut-on transformer une lunette afocale en viseur ? L'image intermédiaire ne peut plus se trouver dans le plan focal image de l'objectif puisque l'objet est à

distance finie. Deux solutions existent :

- ou bien on rend l'objectif plus convergent en ajoutant une lentille convergente supplémentaire qu'on appelle une *bonnette* ;
- ou bien on éloigne l'objectif du plan du réticule où se forme l'image.

En général, un viseur n'est formé que de deux tubes : l'un contenant l'objectif et le réticule et le deuxième l'oculaire. La distance objectif-réticule est donc en général fixée et c'est pour cela qu'on parle de *lunette à frontale fixe* (Cf. figure 12.18). L'oculaire est réglable, pour pouvoir l'adapter à chaque utilisateur.

Le réglage d'un viseur est donc plus simple que celui d'une lunette à l'infini.

Il suffit de régler l'oculaire de manière à voir net le réticule puis de déplacer le viseur sur le banc d'optique jusqu'à voir net l'objet visé.

Ce qu'il est important de retenir, c'est que, puisque la distance objectif-image est constante, la distance objet-viseur est constante elle-aussi.

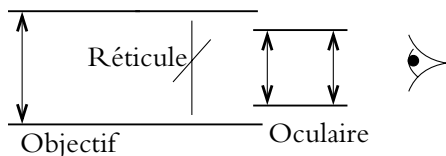


Figure 12.18 Structure d'un viseur.

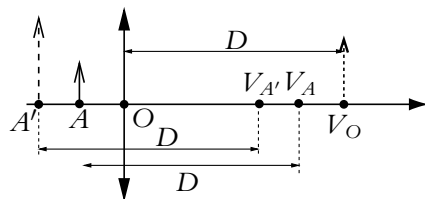


Figure 12.19 Méthode de visée.

Comment faire une mesure relative avec un viseur ?

On étudie le cas d'une lentille de distance focale $f' = 10$ cm et d'un objet placé à 5 cm avant la lentille. On désire vérifier la relation de conjugaison de Descartes :

$$\frac{1}{OA'} - \frac{1}{OA} = \frac{1}{f'}$$

Pour cela, on vise d'abord l'objet A et on repère la position du viseur V_A . On place la lentille en une position qu'on appelle O et on la vise en repérant par exemple des traces à sa surface ou une pointe de crayon que l'on tient à proximité immédiate ; la position du viseur est alors V_O . Enfin, on déplace le viseur jusqu'à voir l'image (placée en A') nette et on appelle $V_{A'}$ la position du viseur. Puisque la distance entre le viseur et ce qu'il vise est constante, comme le montre la figure 12.19, les positions relatives des points V_A , $V_{A'}$ et V_O sont les mêmes que celles des points A , A' et O et **il n'y a pas besoin de connaître la distance de visée D** . Pour vérifier la relation de conjugaison, il suffit de remplacer les points A , A' et O par V_A , $V_{A'}$ et V_O :

$$\frac{1}{V_O V_{A'}} - \frac{1}{V_O V_A} = \frac{1}{f'}$$



Utiliser un viseur est le seul moyen de déterminer la position d'une image virtuelle puisqu'on ne peut pas la recueillir sur un écran.

5. Collimateur

Le collimateur sert à fabriquer un objet à l'infini. Il est formé de deux tubes (Cf. figure 12.20) :

- l'un contenant une lentille convergente ;
- l'autre contenant un objet et une ampoule ou seulement une fente que l'on peut éclairer avec une source extérieure.

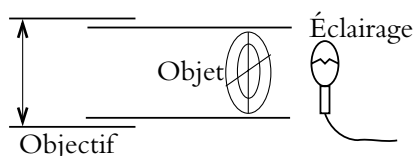


Figure 12.20 Structure d'un collimateur.

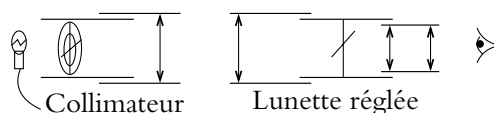


Figure 12.21 Réglage d'un collimateur.

Pour régler le collimateur, il suffit d'amener l'objet ou la fente objet dans le plan focal objet de la lentille en faisant coulisser l'un des tubes par rapport à l'autre. Mais comment savoir si le collimateur est bien réglé ?

Le réglage à l'œil n'est pas possible car l'œil accommode et verra l'image nette même si elle n'est pas à l'infini. **Pour régler un collimateur, on le vise avec une lunette réglée à l'infini et on tourne la bague de réglage du collimateur jusqu'à voir nette l'image** (Cf. figure 12.21). Comme la lunette ne permet d'observer que les objets à l'infini, cela signifie alors que l'image donnée par le collimateur et servant d'objet à la lunette est à l'infini.



Attention, pour éviter l'erreur de parallaxe, il ne faut pas oublier le hochement de tête.

6. Le goniomètre

6.1 Description

Un goniomètre est un appareil de précision qui sert à mesurer des angles et donc des déviations de rayons lumineux par un prisme ou un réseau (Cf. cours de deuxième année). Le programme de première année limite l'étude du goniomètre au prisme.

Un exemple d'appareil est représenté sur la figure 12.22.

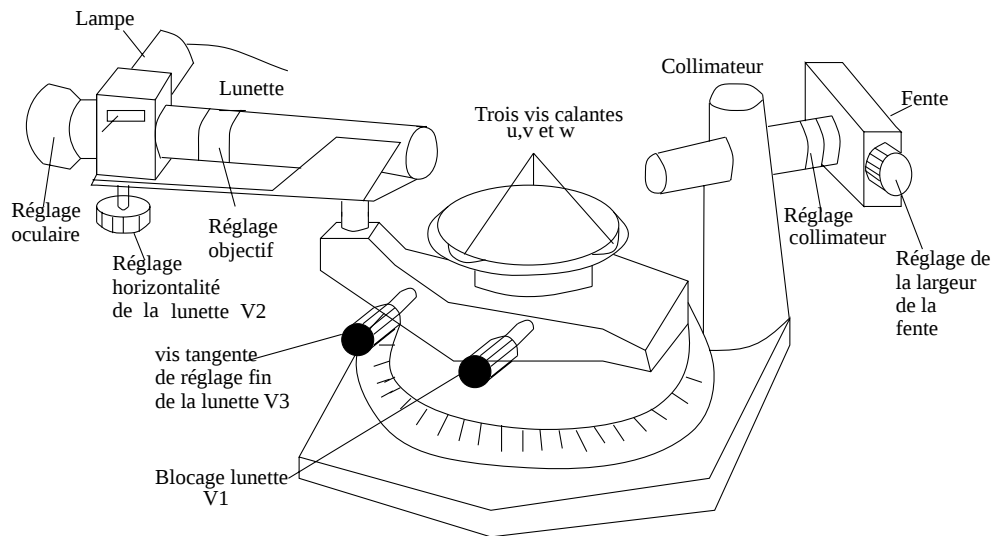


Figure 12.22 Goniomètre.

Il comprend :

- un collimateur dont l'objet est une fente qu'il faut éclairer avec une lampe spectrale ;
- une lunette autocollimatrice montée sur un support mobile autour d'un axe central ;
- un plateau, lui aussi mobile autour de l'axe central.

On peut bloquer le support de la lunette par serrage de la vis V_1 . Cette vis V_1 étant serrée, de légères rotations sont encore possibles au moyen d'une vis tangente V_3 qu'on doit laisser au milieu de sa course et n'utiliser que pour parfaire les pointés. Il existe aussi une vis de blocage du plateau et une vis de réglage fin qui n'ont pas été représentées pour ne pas alourdir le schéma.

La lunette peut basculer par rotation autour d'un axe horizontal, passant près de l'objectif, au moyen d'une vis V_2 .

Le collimateur est supporté par une potence fixe. La fente d'ouverture réglable peut être déplacée par rapport à l'objectif par rotation de la bague de réglage du collimateur.

Le plateau comprend deux plates-formes solidaire, tournant autour de l'axe central. La plate-forme supérieure est mobile verticalement par l'intermédiaire de 3 vis calantes u, v et w.

6.2 Réglage

Puisqu'on cherche à effectuer des mesures précises, il est nécessaire de régler parfaitement l'appareil grâce à la démarche donnée ici. Le réglage s'effectue en quatre étapes.

Les réglages décrits dans les paragraphes c) et d) ne sont pas des connaissances exigibles. Cependant, comme de nombreux goniomètres ne sont pas préréglés et nécessitent ces réglages, on les donne à titre d'information.

a) Lunette autocollimatrice

C'est le premier instrument à régler. On rappelle la procédure mais, pour plus de détails, on se reportera au paragraphe 4.3.

BUT : la lunette doit donner d'un objet à l'infini une image au P.R. de l'œil.

1. Éclairer le réticule en allumant l'ampoule et en baissant la lame semi-réfléchissante. Pour vérifier que la lame est abaissée, il suffit de vérifier que de la lumière sort par l'objectif.
2. Régler l'oculaire.
3. Régler l'objectif par autocollimation soit sur un miroir plan, soit sur une lame à faces parallèles ou une face du prisme posé sur le plateau (il peut être nécessaire d'agir sur les vis calantes du plateau pour que la lumière réfléchie repasse par l'objectif).
4. Vérifier qu'il n'y a pas de parallaxe en hochant la tête.

b) Collimateur

BUT : le collimateur doit donner de la fente une image à l'infini (voir paragraphe 5).

L'éclairer à l'aide d'une lampe spectrale, choisir une fente assez large et le viser avec la lunette déjà réglée à l'infini. Tourner la bague du collimateur jusqu'à voir la fente nette, surtout sur les bords. Une fois le réglage effectué, rendre la fente fine pour les réglages ultérieurs.

c) Réglage de l'axe de la lunette perpendiculairement à l'axe central

BUT : il s'agit de placer la lunette de façon à ce que, par rotation autour de l'axe central, l'axe de la lunette balaie un plan et non un cône. Sinon la mesure des angles de visée est fautive. Pour cela, il est nécessaire que l'axe optique de la lunette soit parfaitement perpendiculaire à l'axe de rotation. Le réglage se fait par autocollimation sur une lame à faces parallèles.

► **Remarque** : Ce réglage n'existe pas sur tous les goniomètres.

C'est la rotation possible de la lunette autour d'un axe horizontal qui permet ce réglage.

1. On rend d'abord la plate-forme grossièrement horizontale (en fait perpendiculaire à l'axe de rotation). Pour cela, descendre le plateau complètement puis le remonter en tournant chaque vis calante (u,v,w) de quelques tours (entre 3 et 10 suivant les goniomètres).

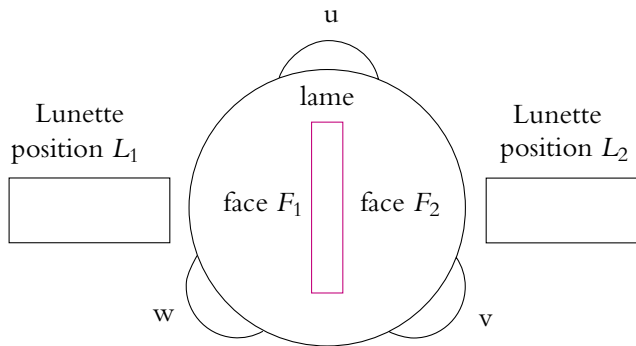


Figure 12.23 Positionnement de la lame.

2. On pose la lame de verre à faces parallèles au centre du plateau de façon à ce que sa base d'appui soit confondue avec une hauteur du triangle équilatéral formé par les trois vis calantes (Cf. figure 12.23). Il faut rendre une des vis inopérante (u dans le cas de la figure) car le réglage avec trois vis serait difficile.
3. On règle la vis V_2 (Cf. figure 12.22) de la lunette à mi-course.
4. On tourne la lunette pour avoir autocollimation sur la face F_1 de la lame (position L_1 sur la figure 12.23). On doit observer une portion de disque lumineux avec l'image du réticule (Cf. figure 12.24) sinon il faut vérifier les précédents points 1, 2 et 3 du réglage. On centre les figures de manière à ce que les traits verticaux des réticules se superposent. Le but est de faire coïncider l'image du réticule $\mathcal{R}'$ avec le réticule $\mathcal{R}$; il reste donc à superposer les traits horizontaux.
Pour cela, agir sur la vis V_2 (horizontalité de la lunette) pour rattraper la moitié du décalage (Cf. figure 12.25) et sur la vis calante (w) la plus proche pour rattraper l'autre moitié du décalage (Cf. figure 12.26). La lumière issue de $\mathcal{R}$ revient alors exactement en $\mathcal{R}'$: $\mathcal{R}'$ est confondu avec $\mathcal{R}$, ce qui indique que la lunette est perpendiculaire à F_1 (Cf. figure 12.27). Mais en général, les faces de F_1 ne sont pas parallèles à l'axe de rotation du plateau. Elle sont inclinées de ϵ par rapport à l'axe et donc l'axe de la lunette n'est pas perpendiculaire à l'axe de rotation.
5. On tourne la lunette de 180° (ou, ce qui revient au même, on tourne le plateau) de manière à amener la deuxième face F_2 devant la lunette en position L_2 sur la figure 12.23.

L'incidence des rayons sur F_2 est alors 2ϵ comme sur la figure 12.28. Les traits horizontaux de $\mathcal{R}$ et $\mathcal{R}'$ ne sont plus confondus (Cf. figure 12.29). On procède alors comme dans la position 1. On superpose les traits verticaux de $\mathcal{R}$ et $\mathcal{R}'$. Puis on agit sur la vis V_2 de la lunette pour rattraper une moitié du décalage horizontal de $\mathcal{R}$ et $\mathcal{R}'$ puis sur la vis calante la plus proche qui est alors (v) pour rattraper l'autre moitié du décalage.

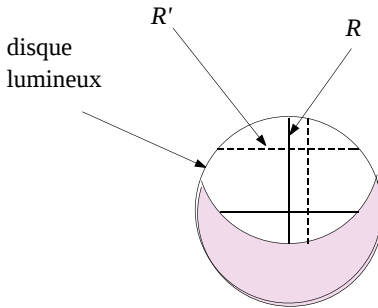


Figure 12.24 Vision initiale.

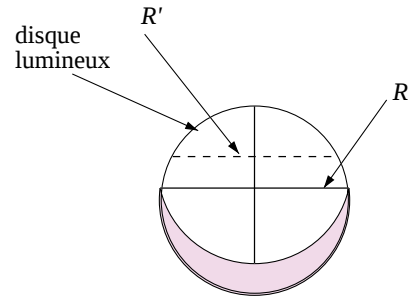


Figure 12.25 Rattrapage de moitié par V_2 .

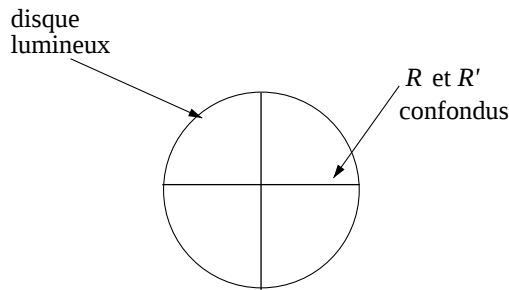


Figure 12.26 Rattrapage de moitié par w .

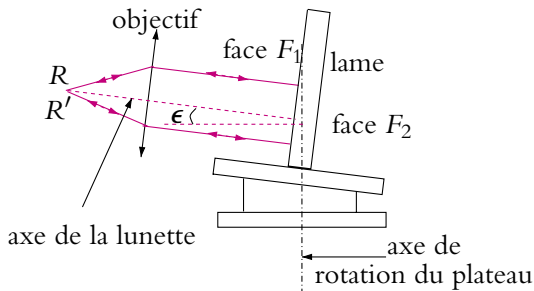


Figure 12.27 Réglage de l'orthogonalité des faces de la lame et des rayons.

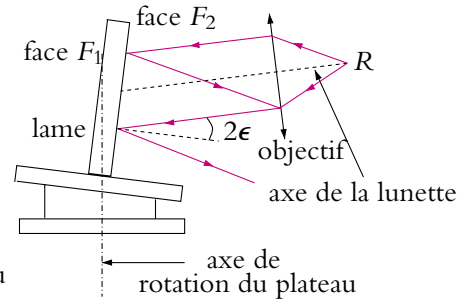


Figure 12.28 Réglage de l'orthogonalité après rotation de 180° .

6. On repasse à l'autocollimation sur F_1 (position L_1) par rotation de 180° de la lunette ou du plateau. On devrait retrouver R et R' en coïncidence exacte, mais il faut en général affiner le réglage car le rattrapage à vue de la moitié de l'écart a été seulement approché et il influe sur les autres opérations. On recommence alors sur F_1 les manœuvres effectuées au point 4 du mode opératoire : rattrapage de l'écart moitié par la vis calante la plus proche (w) et moitié avec V_2 . Puis on

repassé à F_2 (position L_2) etc. jusqu'à coïncidence parfaite de $\mathcal{R}$ et $\mathcal{R}'$ sur F_1 et F_2 (Cf. figure 12.30). Quatre ou cinq demi-tours suffisent en général pour avoir un bon réglage.

L'axe de la lunette est alors perpendiculaire à l'axe de rotation du plateau. **On ne touchera plus à V_2 par la suite.**

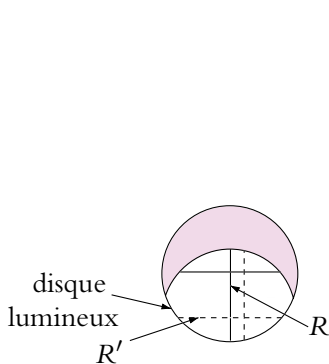


Figure 12.29 Mauvais réglage de l'orthogonalité.

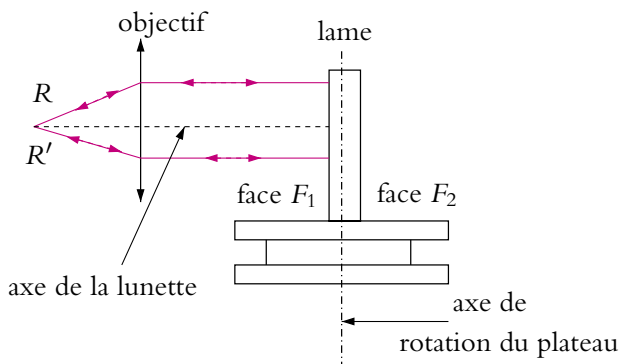


Figure 12.30 Réglage correct de l'orthogonalité.

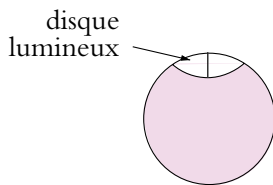


Figure 12.31 Fort décalage des deux positions.

► **Remarque :** Il se peut qu'initialement on ne voit qu'une très faible partie du disque lumineux dans la lunette. Le rattrapage total de l'écart horizontal en position (1) entraînera un décalage trop important en position (2) et on ne verra plus rien dans la lunette (Cf. figure 12.31). Il est préférable, dans ce cas, d'essayer de recentrer petit à petit le disque lumineux avec la vis V_2 .

d) Positionnement du prisme sur la plate-forme

BUT : il s'agit de rendre l'arête utile du prisme parallèle à l'axe de rotation et donc les deux faces utiles du prisme orthogonales aux axes optiques du collimateur et de la lunette (le prisme n'est pas taillé de façon idéale et sa base n'est pas forcément perpendiculaire à l'arête donc il faut incliner la plate-forme pour compenser ce défaut).

1. Rendre de nouveau la plate-forme grossièrement horizontale comme au point (1) du réglage précédent.

2. Repérer l'arête d'angle A utilisée.
3. Imaginer le triangle équilatéral formé par les vis calantes (u , v et w). Placer le prisme comme indiqué sur le schéma de la figure 12.32. L'arête A dépasse légèrement du centre de la plate-forme, les faces utiles F_1 et F_2 sont perpendiculaires aux côtés uw et vw du triangle équilatéral défini à partir des trois vis calantes.

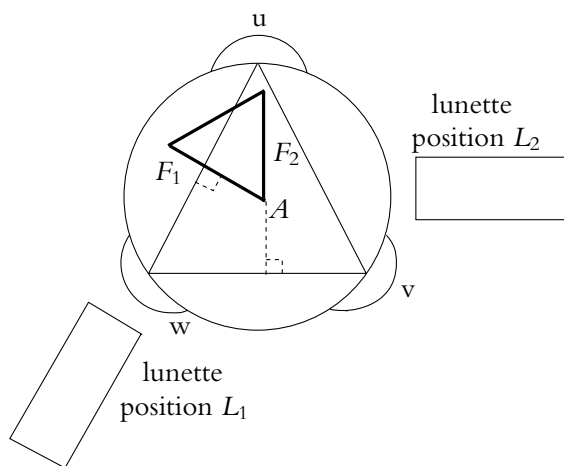


Figure 12.32 Positionnement du prisme.

4. Tourner le plateau pour amener l'arête A à l'opposé du collimateur. Le bloquer. Tourner la lunette pour observer l'autocollimation sur F_1 (position L_1). Amener $\mathcal{R}'$ en coïncidence avec $\mathcal{R}$ en agissant sur la vis tangente de la lunette V_3 pour faire coïncider les traits verticaux et sur la vis calante (w) du plateau pour faire coïncider les traits horizontaux.
5. Tourner ensuite la lunette pour observer l'autocollimation sur F_2 (position L_2). Amener $\mathcal{R}'$ en coïncidence avec $\mathcal{R}$ en agissant sur la vis tangente de la lunette V_3 et sur la vis calante (v) du plateau.
6. Revenir à la position L_1 , réaliser à nouveau la coïncidence, revenir à L_2 , recommencer, etc. jusqu'à coïncidence parfaite.

La face F_1 est alors perpendiculaire à L_1 , la face F_2 est perpendiculaire à L_2 et donc l'arête utile est parallèle à l'axe central car perpendiculaire au plan L_1L_2 (en effet le plan de balayage de l'axe de la lunette est perpendiculaire à l'axe central par le réglage précédent).

Les réglages sont terminés : il ne faut plus déplacer le prisme sur la plate-forme.

► **Remarque :** Pour effectuer des mesures sur les spectres donnés par le prisme, il est nécessaire de relever la lame semi-réfléchissante de la lunette, sinon il y a des images parasites.

6.3 Mesure de l'angle d'un prisme

La première des mesures à effectuer lorsqu'on a réglé le goniomètre est celle de l'angle utile du prisme A .

Il existe deux méthodes.

a) Première méthode

On mesure l'angle entre les deux normales des faces F_1 et F_2 du prisme par autocollimation, c'est-à-dire entre les positions L_1 et L_2 de la lunette sur la figure 12.32. L'angle mesuré entre les deux positions est $\pi - A$, on en déduit la valeur de A (Cf. figure 12.33).

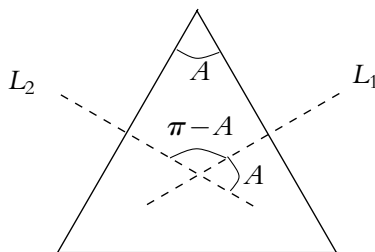


Figure 12.33 Première méthode de mesure de l'angle du prisme.

b) Deuxième méthode

On utilise le collimateur éclairé par une lampe spectrale, avec une fente fine. On oriente le plateau jusqu'à ce que le faisceau éclaire les deux faces utiles du prisme. On vise alors avec la lunette les positions des rayons réfléchis (Cf. figure 12.34). L'angle entre les deux positions est $\beta = 2A$. En effet, d'après la figure 12.34, $\beta = A + x + y$ et $A' = \left(\frac{\pi}{2} - x\right) + \left(\frac{\pi}{2} - y\right)$. Or dans le chapitre sur le spectroscope à prisme, on montrera que $A' = \pi - A$. On en déduit que $x + y = A$ et $\beta = 2A$. Cette deuxième méthode est plus précise que la première car, comme on mesure $2A$ avec la même incertitude Δ que dans la première méthode (2 fois l'incertitude sur la lecture de la position de la lunette), l'incertitude sur A est $\frac{\Delta}{2}$.

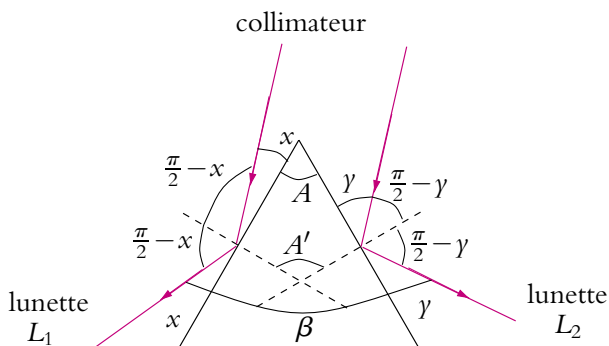


Figure 12.34 Deuxième méthode de mesure de l'angle du prisme.

D'autres exemples d'utilisation du goniomètre seront décrites dans le chapitre « Spectroscopie à prisme ».

A. Applications directes du cours

1. Problème de myopie pour la vision dans un miroir

Un observateur myope a son punctum proximum PP à 15 cm et son punctum remotum PR à 40 cm. Il tient un miroir situé au maximum à 70 cm de son œil. Quel rayon peut avoir le miroir pour qu'il puisse voir l'image d'un objet situé à l'infini ?

2. Système achromatique

On caractérise la dispersion d'un matériau à l'aide de trois longueurs d'onde : la raie rouge de l'hydrogène notée raie C ($\lambda = 656$ nm), la raie jaune du sodium notée raie D ($\lambda = 589$ nm) et la raie bleu-vert de l'hydrogène notée raie F ($\lambda = 486$ nm). Le pouvoir dispersif ν du matériau est donné par la relation

$$\nu = \frac{n_D - 1}{n_F - n_C}$$

où n_C , n_D et n_F sont les indices du matériau pour les raies C, D et F. La vergence V d'une lentille dépend de l'indice par la relation : $V = (n - 1)A$ où A est uniquement une caractéristique géométrique de la lentille.

On cherche à obtenir un système achromatique à l'aide de deux lentilles.

1. On accole deux lentilles de vergence V_1 et V_2 . Etablir que le système est équivalent à une lentille dont on précisera les caractéristiques.
2. Une lentille possède une distance focale $f'_D = 30$ cm et son pouvoir dispersif vaut : $\nu = 50$. Calculer l'écart $f'_C - f'_F$. On justifiera les éventuelles approximations.
3. Quelles sont les conséquences de cet écart ?
4. A quelle condition les foyers des raies C et F sont-ils confondus ?
5. A-t-on une lentille achromatique ? Justifier.
6. On souhaite une lentille de focale moyenne 50 cm avec deux verres caractérisés par $\nu_1 = 30$ et $\nu_2 = 60$. Quelles doivent être les focales des lentilles à utiliser ?

B. Problèmes

1. Microscope d'après Veto 2000 et Centrale PC 2000

Les deux parties sont très largement indépendantes.

On ne fera pas d'applications numériques, les valeurs numériques ne sont données que pour aider à répondre à certaines questions.

1. Etude d'un détecteur optique : l'œil :

On peut modéliser l'œil par une lentille mince convergente $\mathcal{L}$ de centre O et de distance focale variable f' . Le cristallin donne d'un objet une image renversée sur la rétine située à la distance $\overline{OE} = d$ et servant d'écran. On suppose d fixe pour un œil donné.

Un œil au repos voit les objets situés au punctum remotum PR à une distance D_m de O . Quand il accommode au maximum, l'œil voit les objets situés au punctum proximum PP à une distance d_m de O . Pour un œil normal, on a D_m infinie et $d_m = \delta = -25,0$ cm.

A chaque état d'accommodation est associée une distance focale f' pour la lentille modélisant l'œil.

L'œil qu'on étudie vérifie $d = 15,0$ mm, $d_m = -32,3$ cm et $D_m = 1,11$ m.

- a) Placer sur un schéma la lentille, l'écran et les points PP et PR . Hâchurer la zone vue par l'œil.
- b) Quelle est la nature des objets que cet œil peut voir ?
- c) Quel défaut de vision modélise-t-on ici ?
- d) Déterminer la distance focale de la lentille modélisant cet œil au repos.
- e) Même question lorsqu'il accommode au maximum.
- f) Etablir que l'association de deux lentilles minces $\mathcal{L}_1$ et $\mathcal{L}_2$ de distance focale respectivement f'_1 et f'_2 accolées en O est une lentille mince dont on donnera le centre et la distance focale.
- g) On veut corriger la vision de cet œil avec des verres de contact ou lentilles dont on admettra qu'ils sont accolés à la lentille modélisant l'œil. Donner la ou les distances focales des verres de contact pour obtenir les limites de vision d'un œil normal. On précisera le caractère convergent ou divergent de ces verres de contact.

2. Etude géométrique d'un microscope :

Un microscope optique permet d'observer des globules sanguins.

Il est modélisable par deux lentilles minces convergentes $\mathcal{L}_1$ pour l'objectif de distance focale f'_1 et $\mathcal{L}_2$ pour l'oculaire de distance focale f'_2 . Il est réglé pour donner une image à l'infini d'un objet réel AB perpendiculaire à l'axe optique, A étant sur l'axe optique) légèrement en avant du foyer objet de l'objectif. Cette image est observée par un œil emmétrope (normal) placé au voisinage du foyer image de l'oculaire. On notera $A'B'$ l'image intermédiaire.

Le microscope porte les indications suivantes :

- $\times 40$ pour l'objectif, ce qui signifie que la valeur absolue du grandissement de l'objet AB par l'objectif est de 40 ;
- $\times 10$ pour l'oculaire, ce qui signifie que le grossissement commercial ou rapport entre l'angle sous lequel on voit l'image à l'infini d'un objet à travers l'oculaire seul et l'angle sous lequel on voit ce même objet à l'œil nu lorsqu'il est situé à la dimension minimale de vision distincte δ vaut 10 ;
- $\omega_0 = 0,65$ pour l'ouverture numérique ou valeur de $n \sin u$ avec n le milieu dans lequel se trouve l'objectif et u l'angle maximum des rayons issus de A arrivant sur l'objectif.
- $\Delta = 16$ cm pour l'intervalle optique ou distance entre le foyer image F'_1 de l'objectif et le foyer objet F_2 de l'oculaire.

a) Faire un schéma du dispositif (sans respecter l'échelle) et tracer la marche de deux rayons lumineux issus du point B de l'objet AB , l'un émis parallèlement à l'axe optique et l'autre passant par le foyer objet de l'objectif.

b) En utilisant le grossissement commercial, déterminer la distance focale f'_2 de l'oculaire.

- c) Déterminer la distance focale f'_1 de l'objectif. On pourra utiliser le grandissement de l'objectif.
- d) Calculer la distance O_1A permettant de positionner l'objet.
- e) Déterminer la latitude de mise au point à savoir la variation de la distance O_1A compatible avec l'observation d'une image par l'œil situé au foyer image de l'oculaire.
- f) Calculer le grossissement commercial pour une image finale à l'infini.
- g) Calculer l'angle u intervenant dans l'ouverture numérique pour un objectif placé dans l'air. Le microscope est-il utilisé dans les conditions de Gauss ? Quel type d'aberrations doit-on corriger ? Quel est l'ordre de grandeur du diamètre de la monture de l'objectif ?
- h) Déterminer la position et la taille du cercle oculaire défini comme l'image de la monture de l'objectif à travers l'oculaire. Quel est l'intérêt de placer l'œil dans le plan du cercle oculaire ?

2. Lunette de Galilée d'après CCP MP 2007

Une lunette de Galilée comprend un objectif assimilable à une lentille mince $\mathcal{L}_1$ de centre O_1 et de vergence $V_1 = 5,0$ dioptries et un oculaire assimilable à une lentille mince $\mathcal{L}_2$ de centre O_2 et de vergence $V_2 = -20$ dioptries.

1. Déterminer la nature des deux lentilles et donner la valeur de leur distance focale f'_1 et f'_2 .
2. La lunette est de type afocale. Préciser dans ces conditions la position relative des deux lentilles en donnant la valeur de $d = \overline{O_1O_2}$.
3. Tracer, dans les conditions de Gauss, la marche d'un rayon lumineux incident arrivant d'un objet à l'infini, faisant un angle θ avec l'axe optique et émergeant sous l'angle θ' .
4. En déduire le grossissement ou grandissement angulaire en fonction de θ et θ' puis des distances focales des deux lentilles. Donner sa valeur numérique.
5. Un astronome amateur utilise cette lunette normalement adaptée à la vision d'objets terrestres pour observer deux cratères lunaires : Copernic de diamètre 96 km et Clavius de diamètre 240 km. On rappelle que la distance entre la Terre et la Lune vaut 384.10^6 m. Voit-il ces deux cratères à l'œil nu ? à l'aide de la lunette ?
6. Déterminer le grandissement du dispositif.

13

T.P. : Focométrie

On va présenter différentes méthodes de mesure de distances focales de miroirs et de lentilles. On encourage le lecteur à refaire les constructions pour chaque méthode présentée. On donne l'évaluation de l'incertitude associée.

1. Reconnaissance rapide du caractère d'une lentille ou d'un miroir

Avant toute mesure, il est nécessaire de savoir à quel type de système on a affaire :

- lentille convergente ou divergente ;
- miroir plan, concave ou convexe.

Certaines règles permettent de déterminer rapidement de quel système il s'agit.

1.1 Miroirs

La première chose à faire est d'observer la surface du miroir, il est alors possible de déterminer sa forme. Sinon, il faut se regarder dans le miroir en commençant près du miroir puis en reculant.

a) Miroir plan

Chacun a l'habitude de se regarder dans un miroir plan. L'image est droite et de même grandeur.

b) Miroir concave

On commence par se regarder en étant proche du miroir, puis on recule. Les trois positions possibles sont récapitulées sur les figures 13.1 à 13.3. On se trouve à la place de l'objet AB .

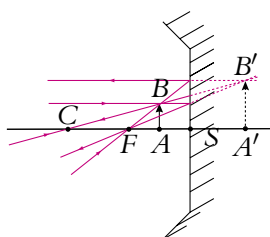


Figure 13.1 Image d'un objet entre le foyer et le sommet.

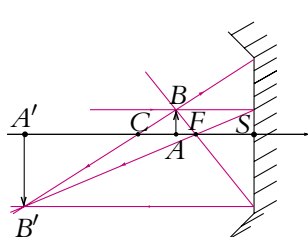


Figure 13.2 Image d'un objet entre le centre et le foyer.

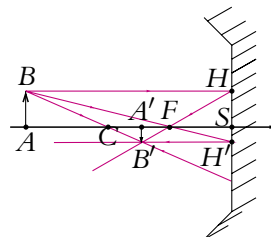


Figure 13.3 Image d'un objet avant le centre.

1. Si l'on est près du miroir (entre F et S), l'image est virtuelle, droite et plus grande (figure 13.1). Le miroir est grossissant, on se voit donc dans le miroir, plus grand. C'est le cas des petits miroirs de maquillage.
2. Si l'on se trouve entre C et F , l'image est réelle renversée derrière C (figure 13.2). On ne peut observer l'image qui est derrière soi. On ne voit qu'une image floue.
3. Lorsqu'on se trouve avant C , l'image est réelle, renversée et plus petite devant soi (figure 13.3). On se voit donc plus petit et à l'envers.

c) Miroir convexe

Dans ce cas, une seule position d'objet réel est possible, puisque C et F sont virtuels.

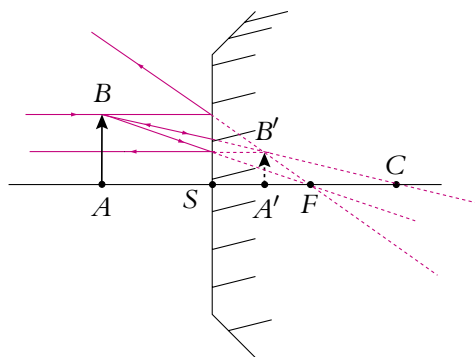


Figure 13.4 Miroir convexe

On se voit droit, plus petit et dans le miroir (figure 13.4).

1.2 Lentilles

Contrairement aux miroirs, on ne peut être à la fois l'objet et l'observateur par une lentille. On observe alors d'abord un objet proche puis un objet lointain.

a) Lentilles convergentes

Les deux cas d'objets réels sont redonnées sur les figures 13.5 et 13.6.

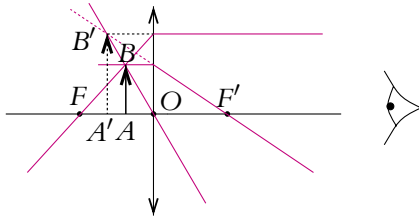


Figure 13.5 Image d'un objet proche.

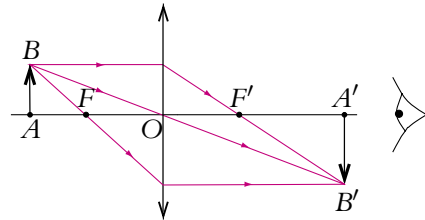


Figure 13.6 Image d'un objet lointain.

Si l'objet est proche, la lentille agit comme une loupe : on voit donc l'objet plus gros et virtuel.

Si l'objet est lointain, on voit une image renversée. Si l'objet est suffisamment lointain, l'image se forme dans le plan focal image. Pour la voir nette, il est nécessaire qu'elle soit au-delà du P.P. de l'œil. Il est préférable de tenir la lentille à bout de bras, surtout dans le cas de grande distance focale. Dans le cas contraire, on verra une image renversée floue.

b) Lentille divergente

Il n'y a qu'une seule possibilité : un objet réel (proche ou lointain) donne une image virtuelle droite plus petite (figure 13.7).

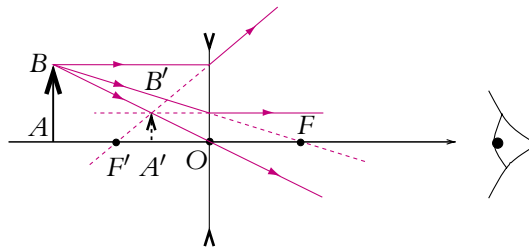


Figure 13.7 Image d'un objet réel par une lentille divergente.

2. Méthode directe

2.1 Lentille convergente

a) Principe

On utilise le fait que l'image par une lentille d'un objet à l'infini est dans le plan focal image.

b) Manipulation

On utilise un collimateur, réglé à l'aide d'une lunette de visée autocollimatrice. On place la lentille et on déplace un écran jusqu'à obtenir l'image du réticule du collimateur nette : on en déduit la position de F' et la distance focale de la lentille $f' = OF'$ qui est égale à la distance ℓ entre la lentille et l'écran. On peut aussi utiliser un viseur et viser successivement la lentille puis l'image du réticule.

c) Calculs d'incertitude

Puisqu'on repère deux positions il faut additionner les deux erreurs de lectures : incertitude sur la position de la lentille $\Delta\ell_1$ et incertitude sur la position de l'image $\Delta\ell_2$. On prend en général pour les incertitudes de lecture une demi-graduation (0,5 mm) pour un banc d'optique. Il faut estimer aussi l'incertitude due à la latitude de mise au point. Pour cela, il faut déterminer l'intervalle des positions $\Delta\ell_3$ de l'écran ou du viseur pour lesquelles on estime que l'image est nette.

L'incertitude totale est :

$$\Delta f' = \Delta\ell = \Delta\ell_1 + \Delta\ell_2 + \Delta\ell_3$$

2.2 Lentille divergente

a) Principe

La méthode n'est pas applicable directement (foyers virtuels). On doit donc accoler à la lentille divergente (vergence V_d) une lentille convergente de plus grande vergence absolue V_c , de manière à réaliser un système convergent. On a montré dans le chapitre sur les lentilles que la vergence V de l'ensemble formé par deux lentilles accolées est la somme des vergences des deux lentilles :

$$V = V_c + V_d$$

b) Manipulation

On applique la même méthode qu'avec la lentille convergente pour déterminer V et on en déduit $V_d = V - V_c$.

c) Détermination de l'incertitude

On cherche l'incertitude sur la vergence de la lentille divergente V_d . Soit :

$$V_d = V - V_c \Rightarrow \Delta V_d = \Delta V + \Delta V_c$$

Mais avec cette méthode, on mesure l'incertitude sur la distance focale totale $\Delta\ell = \Delta f'$. En utilisant les différentielles logarithmiques (on rappelle que la vergence est l'inverse de la distance focale), on obtient :

$$\frac{\Delta V}{V} = \frac{\Delta f'}{f'} \Rightarrow \Delta V_d = \Delta V_c + (V_c + V_d) \frac{\Delta\ell}{\ell}$$

ΔV_c est l'incertitude sur la vergence de la lentille convergente.

2.3 Miroir concave

a) Principe

Lorsque l'objet est dans le plan contenant le centre du miroir, l'image est aussi dans ce plan, de même taille que l'objet et renversée.

b) Manipulation

Si on utilise, par exemple, une diapositive comme objet, on déplace le miroir jusqu'à obtenir l'image sur la diapositive. Le rayon du miroir est égal à la distance entre le miroir et la diapositive.

c) Explications

La construction de la figure 13.8 montre que l'image est dans le même plan que l'objet et l'application de la formule de grandissement avec origine au sommet montre que le grandissement vaut -1 .

$$\gamma = \frac{\overline{A'B'}}{\overline{AB}} = -\frac{\overline{SA'}}{\overline{SA}} = -1$$

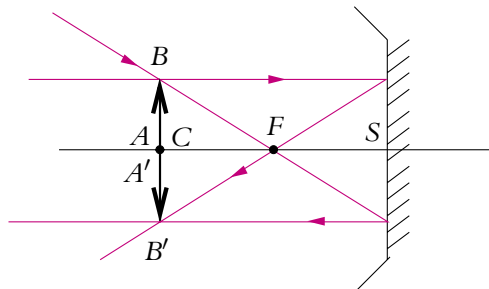


Figure 13.8 Grandissement -1 pour un miroir.

3. Autocollimation

3.1 Lentille convergente

a) Manipulation

On place la lentille convergente et on tient le miroir plan derrière la lentille, à la main en l'inclinant légèrement, ce qui ne change pas grand chose. On cherche la position de la lentille donnant une image nette sur le cache de l'objet et on mesure la distance lentille-objet ℓ qui est égale à la distance focale f' de la lentille. Une fois le réglage effectué, on peut aussi viser la lentille grâce à un viseur, puis la retirer et viser l'objet (plus précis).

b) Explication

On place un miroir plan devant la lentille convergente (cf. figure 13.9).

L'objet A va donner une image A' à travers le système lentille-miroir. On appelle A_1 et A_2 les images intermédiaires :

$$A \xrightarrow{\text{lentille}} A_1 \xrightarrow{\text{miroir}} A_2 \xrightarrow{\text{lentille}} A'$$

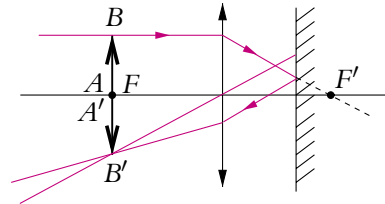


Figure 13.9 Autocollimation



Attention, puisque la lumière change de sens, il faut penser dans les raisonnements à inverser les foyers objet et image, sachant que les deux foyers sont symétriques par rapport à la lentille.

On suppose que A est dans le plan focal objet de la lentille, alors A_1 est à l'infini, de même que A_2 . Le point A' est alors dans le plan focal devenu le plan focal image. Ainsi $A = A'$: l'objet et l'image sont dans le même plan et on les voit **simultanément nets**. D'autre part si l'on note O le centre optique de la lentille, et que l'on dispose d'un objet AB , le grandissement est -1 (image renversée) puisque :

$$\gamma = \frac{\overline{A'B'}}{\overline{AB}} = \frac{\overline{OA'}}{\overline{OA}} = -1$$

Sur la figure 13.9, le point A est pris sur l'axe optique. Si ce point n'est pas sur l'axe, le résultat est identique, l'image est simplement décalée dans le plan focal.

3.2 Lentille divergente

Le problème est le même que pour la méthode directe ; on procède de même en accolant la lentille divergente à une lentille convergente de plus grande vergence (la mise au point est plus délicate).

4. Méthode de Silberman

4.1 Lentille convergente

a) Principe

La distance minimale D entre un objet réel et son image réelle est égale à $4f'$. Dans cette configuration, la lentille est à mi-distance de l'objet et de l'écran.

b) Manipulation

Sachant que pour $D = 4f'$, l'image et l'objet sont symétriques par rapport à la lentille, on place l'objet et l'écran contre la lentille et on les écarte symétriquement jusqu'à voir une image nette. On ajuste ensuite le réglage des positions pour avoir le même écart lentille-écran et lentille-objet, puis on mesure D .

c) Explications

On reprend une démonstration effectuée dans le chapitre sur les lentilles (paragraphe sur les projections d'images). On considère une lentille convergente de distance focale image f' et de centre optique O . Soit un objet A réel, on cherche la condition pour obtenir une image A' réelle. On note $x = \overline{OA'} > 0$ et $D = \overline{AA'} > 0$. La relation de conjugaison de Descartes donne :

$$\frac{1}{\overline{OA'}} - \frac{1}{\overline{OA}} = \frac{1}{f'} \Rightarrow (\overline{OA} - \overline{OA'})f' = \overline{OA} \overline{OA'}$$

or $\overline{OA} = \overline{OA'} + \overline{A'A}$, soit :

$$x^2 - xD + Df' = 0$$

On calcule le discriminant qui doit être positif puisqu'on cherche des solutions réelles :

$$\Delta = D^2 - 4f'D = D(D - 4f') \geq 0 \Leftrightarrow D \geq 4f' \text{ car } D > 0$$

Pour avoir une image réelle, il faut donc que $D \geq 4f'$. Dans le cas particulier où $D = 4f'$, $\Delta = 0$ et il y a une seule solution à l'équation, donc une seule image telle que $x = \overline{OA'} = \frac{D}{2}$. Ainsi la lentille est à mi-chemin entre l'objet et l'image.

4.2 Lentille divergente

On procède de la même manière en l'accolant à une lentille convergente de plus grande vergence absolue : on détermine la distance focale équivalente à l'association et on en déduit f pour la lentille divergente.

5. Méthode de Bessel

5.1 Lentille convergente

a) Principe

L'objet et l'écran sont à la distance $D > 4f'$ fixes l'un par rapport à l'autre. Il existe alors deux positions de la lentille O_1 et O_2 telles que l'objet et l'écran soient conjugués. Si la distance $O_1 O_2$ est notée d , on obtient :

$$f' = \frac{D^2 - d^2}{4D}$$

b) Manipulation

On repère les 2 positions du support de la lentille correspondant à une image sur l'écran et on mesure d . On recommence la mesure pour plusieurs valeurs de D et on

trace $(d/D)^2$ en fonction de $1/D$. On cherche alors la droite de meilleure interpolation. On déduit de la valeur de sa pente la distance focale f' . On peut améliorer le résultat en ajoutant le point $(0,1)$ aux points de mesures.

► **Remarque :** Une seule mesure de d et D suffit pour déterminer f' , mais en faire plusieurs permet d'améliorer la précision.

c) Explication

Il suffit de déterminer les solutions de l'équation étudiée dans le paragraphe sur la méthode de Silberman. Elles sont, pour $D > 4f'$:

$$\overline{OA'} = \frac{D}{2} \pm \frac{1}{2} \sqrt{D^2 - 4f'D}$$

Si on garde A' fixe et qu'on déplace la lentille, on obtient alors les deux positions de O_1 et O_2 pour avoir une image en A' :

$$\overline{O_1A'} = \frac{D}{2} - \frac{1}{2} \sqrt{D^2 - 4f'D} \quad \text{et} \quad \overline{O_2A'} = \frac{D}{2} + \frac{1}{2} \sqrt{D^2 - 4f'D}$$

soit en prenant la différence entre les deux équations précédentes :

$$\overline{O_1O_2} = \sqrt{D^2 - 4f'D}$$

On peut alors tracer la droite $\left(\frac{d}{D}\right)^2$ en fonction de $\frac{d}{D}$ qui est de pente $-4f'$.

6. Méthode des points conjugués

6.1 Lentille convergente

a) Principe

C'est une exploitation de la loi de conjugaison de Descartes. On détermine pour diverses positions de l'objet A la position de l'image A' et on trace la droite $\frac{1}{\overline{OA'}}$ en fonction de $\frac{1}{\overline{OA}}$. En écrivant la relation de conjugaison de Descartes :

$$\frac{1}{\overline{OA'}} = \frac{1}{\overline{OA}} + \frac{1}{f'}$$

on trouve que $\frac{1}{f'}$ est l'ordonnée à l'origine de la droite tracée.

b) Manipulation

On dispose une lentille sur le banc d'optique et on repère sa position. Pour chaque position de l'objet, on cherche la position de l'image (réelle); pour une meilleure précision on pourra faire des pointés avec le viseur (on vise l'image, puis on retire la lentille en repérant la position du support et on vise ensuite l'objet). On trace le graphe représentant $\frac{1}{OA'}$ en fonction de $\frac{1}{OA}$.

6.2 Lentille divergente

On trace la droite $\frac{1}{OA'}$ en fonction de $\frac{1}{OA}$ pour un objet virtuel (obtenu grâce à une lentille convergente auxiliaire) et une image réelle : on effectue les pointés avec un viseur.

T.P. : Spectroscopie à prisme

14

Dans ce chapitre, on présente un exemple d'utilisation du goniomètre : le spectroscopie à prisme.

1. Présentation

Lors de l'étude des sources lumineuses, on a présenté un dispositif à prisme permettant d'observer le spectre d'une source lumineuse. Un spectre lumineux est caractéristique d'un corps si bien que l'étude de ces spectres est très importante car elle permet d'identifier les corps qui sont à l'origine de cette lumière : on a cité précédemment les lampes spectrales mais, en astrophysique par exemple, cela permet de connaître les éléments chimiques constituant une étoile. Le montage présenté dans le chapitre sur les sources permet une étude qualitative mais ne permet pas de déterminer avec précision les longueurs d'onde des raies présentes dans un spectre, ce qui est possible avec les spectroscopes.

En première année, on se limite à l'utilisation des spectroscopes à prisme. On commence par étudier le prisme.

2. Application des lois de Descartes au prisme

2.1 Description

Un prisme est un bloc de verre pyramidal (Cf. figure 14.1). On utilise deux des faces dont l'arête commune est appelée *arête utile*.

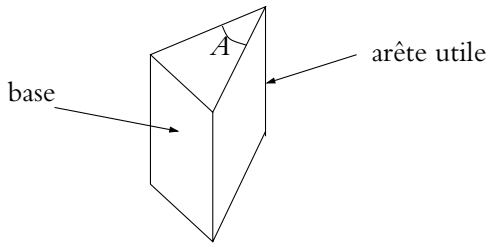


Figure 14.1 Prisme.

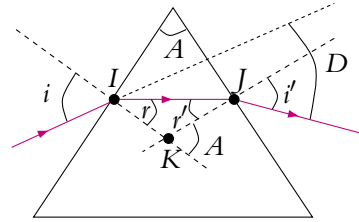


Figure 14.2 Angles du prisme.

L'angle entre les deux faces utiles est notée A . On utilise des verres de différents indices donnés en général pour le jaune ($\lambda = 0,56 \mu\text{m}$), longueur d'onde moyenne du spectre. On distingue en fait deux types de verres : le flint ($n \simeq 1,7$) et le crown ($n \simeq 1,5$) sachant que certains fabricants proposent maintenant des prismes en crown ayant un indice proche de 1,7. L'indice de réfraction du verre dépend de la longueur d'onde selon la loi de Cauchy :

$$n = a + \frac{b}{\lambda^2} \quad (14.1)$$

a et b étant des constantes positives qui dépendent du verre. On dit que le prisme est dispersif, c'est-à-dire qu'il dévie différemment des lumières de différentes longueurs d'onde : il sert à étudier les spectres des sources de lumière.

2.2 Relations du prisme

Il s'agit d'établir les relations entre les différents angles intervenant dans les constructions relatives au prisme. Les notations sont celles de la figure 14.2 réalisée dans le plan d'incidence. Par habitude, sur le prisme, on raisonne avec des angles positifs non orientés (tous les angles sont compris entre 0 et $\pi/2$). Le milieu extérieur est l'air d'indice 1 et l'indice de réfraction du verre est noté n .

Aux points I et J , on peut appliquer les relations de Descartes soit :

$$n \sin r = \sin i \quad \text{et} \quad n \sin r' = \sin i'$$

L'angle entre les deux normales est égal à l'angle entre les deux faces utilisées du prisme soit A . Ainsi dans le triangle IJK , la somme des angles donne :

$$r + r' + (\pi - A) = \pi \quad \Leftrightarrow \quad r + r' = A$$

Il reste à déterminer la déviation, c'est-à-dire l'angle D entre le rayon émergent du prisme et le rayon incident. Pour cela, on calcule l'angle dont « tourne » le rayon à chaque interface.

- En I , le rayon tourne d'un angle $\alpha = i - r$.
- En J , le rayon tourne d'un angle $\gamma = i' - r'$.

Globalement, le rayon tourne d'un angle $D = \alpha + \gamma = i + i' - r - r'$, soit avec $r + r' = A$:

$$D = i + i' - A$$

On peut récapituler les quatre formules du prisme :

$$\begin{aligned} n \sin r &= \sin i \\ n \sin r' &= \sin i' \\ r + r' &= A \\ D &= i + i' - A \end{aligned} \quad (14.2)$$

2.3 Dispersion par un prisme

Puisque les angles de déviation dépendent de la longueur d'onde, on va déterminer la couleur qui est la plus déviée. Pour cela, on considère un prisme d'angle A éclairé par une lumière polychromatique sous une incidence i . Ainsi l'angle d'incidence est le même quelle que soit la longueur d'onde. D'après une des relations précédentes, la déviation dépend de A constant, de i constant et de i' ; il suffit donc d'étudier la variation de i' en fonction de λ pour déterminer celle de D .

Si λ croît, la relation de Cauchy (11.13) impose à l'indice n de décroître. Or $\sin r = \frac{1}{n} \sin i$ donc r croît. La relation entre r et r' donne $r' = A - r$ donc r' décroît. Enfin $\sin i' = n \sin r'$ donc comme n et r' décroissent, i' décroît.

On en déduit que, si la longueur d'onde croît, i' et $D = i + i' - A$ décroissent. **Le prisme dévie donc plus le bleu que le rouge.**

2.4 Condition d'émergence par la deuxième face

En I , la lumière passe de l'air à un milieu plus réfringent donc il ne peut y avoir de réflexion totale ; par contre, en J , c'est possible. On va déterminer la condition sur l'angle d'incidence pour que la lumière puisse être réfractée sur la deuxième face et non réfléchie.

En J , la lumière passe d'un milieu d'indice n à un milieu d'indice 1. Pour qu'il n'y ait pas réflexion totale, il faut que $r' < R_{\text{lim}}$ avec $\sin R_{\text{lim}} = \frac{1}{n}$, donné par la relation (9.3) avec $n_1 = n$ et $n_2 = 1$. Ainsi il n'y a pas réflexion totale en J si :

$$r' \leq \text{Arcsin} \left(\frac{1}{n} \right) \quad (14.3)$$

D'autre part, le rayon pénètre dans le prisme en I donc $\sin i = n \sin r$ soit, puisque $\sin i \leq 1$, $\sin r = \frac{1}{n} \sin i \leq \frac{1}{n}$. On en déduit :

$$r \leq R_{\text{lim}} \quad (14.4)$$

Avec la relation $A = r + r'$ et les deux relations précédentes, on peut écrire une première condition pour qu'il n'y ait pas réflexion totale sur la deuxième face :

$$A \leq 2R_{\text{lim}} \quad (14.5)$$

Par exemple, pour $n = 1,5$, $R_{\text{lim}} = 41,8^\circ$. Ainsi, un prisme à angle droit (Cf. figure 14.3) sera un prisme à réflexion totale. On utilise des prismes à réflexion totale dans les appareils photos à « visée reflex », c'est-à-dire à visée à travers l'objectif. On se sert de prismes plutôt que de miroirs car ils sont plus faciles à manipuler et moins fragiles (en particulier le verre ne se corrode pas).

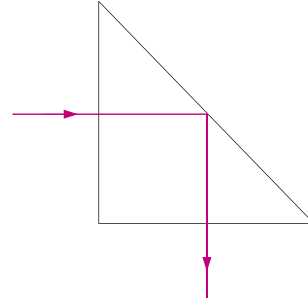


Figure 14.3 Prisme à réflexion totale.

Quelle est maintenant la condition sur l'angle d'incidence i ?

À la limite, $r' = R_{\text{lim}}$ donc $r = A - R_{\text{lim}}$. On peut calculer alors l'angle limite d'incidence i_{lim} par :

$$\sin i_{\text{lim}} = n \sin r = n \sin(A - R_{\text{lim}}) \quad (14.6)$$

Si $i > i_{\text{lim}}$, $r > A - R_{\text{lim}}$ et $r' < R_{\text{lim}}$: il n'y a pas réflexion totale. Pour $A = 60^\circ$ et $n = 1,5$, on trouve $i_{\text{lim}} = 27,9^\circ$.

Pour ne pas avoir de réflexion totale, il est nécessaire que $i_{\text{lim}} < i < \frac{\pi}{2}$ donc on a toujours intérêt à envoyer de la lumière sur un prisme avec une forte incidence.

2.5 Existence d'un minimum de déviation

Si on éclaire le prisme en lumière monochromatique et si on fait varier l'angle d'incidence i en partant de $i = \frac{\pi}{2}$, on s'aperçoit expérimentalement que la déviation diminue, passe par un minimum noté D_m puis augmente à nouveau. Il s'agit de déterminer une relation entre D_m , n et A .

Pour cela, on utilise le retour inverse de la lumière. Un angle d'incidence i donne l'angle d'émergence i' et par retour inverse de la lumière, un angle d'incidence i' donnera l'angle d'émergence i . Étant donné que $D = i + i' - A$, les angles d'incidence i et i' donneront la même déviation. Ainsi, dans le cas général, la déviation est la même pour deux angles donc la forme de la courbe de la déviation D en fonction de l'angle d'incidence est l'une des deux représentées sur la figure 14.4. Elle présente donc

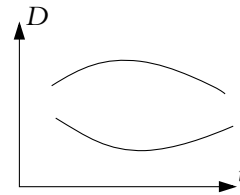


Figure 14.4 Minimum de déviation.

un minimum ou un maximum obtenu d'après ce qui précède lorsque $i = i'$ (retour inverse de la lumière). Expérimentalement, on trouve qu'il s'agit d'un minimum.

Les relations (14.2) donnent avec $i = i'$:

$$r = r' = \frac{A}{2} \quad \text{et} \quad D_m = i + i' - A$$

soit avec $\sin i = n \sin r$:

$$\sin \frac{D_m + A}{2} = n \sin \frac{A}{2} \quad (14.7)$$

Cette relation est utilisée en travaux pratiques pour déterminer l'indice d'un prisme comme cela est décrit dans le paragraphe suivant.

2.6 Mesure de l'indice d'un prisme

Cette étude est réalisée avec un goniomètre.

On désire mesurer l'indice pour une longueur d'onde donnée, par exemple le jaune du sodium. Pour cela, on utilise la relation établie au minimum de déviation :

$$n = \frac{\sin \left(\frac{A + D_m}{2} \right)}{\sin \left(\frac{A}{2} \right)} \quad (14.8)$$

Il faut se rappeler que, d'après l'étude faite sur le prisme, on a intérêt à l'éclairer en incidence rasante pour être sûr de trouver des rayons émergents.

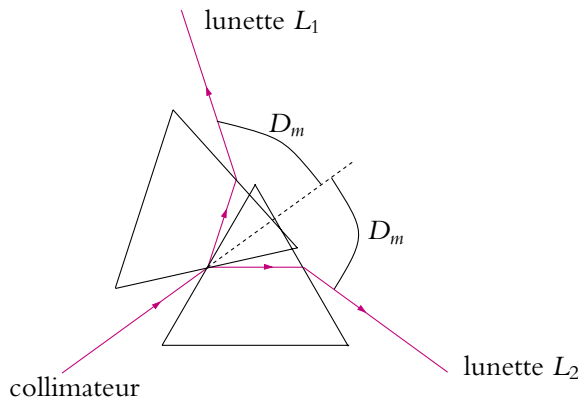


Figure 14.5 Mesure de D_m .

On commence donc par éclairer le prisme par le collimateur avec une forte incidence et on observe à l'œil le rayon réfracté. On tourne le prisme de manière à faire diminuer l'angle d'incidence toujours en observant l'image réfractée. Pour un certain angle, on voit l'image s'arrêter et repartir dans l'autre sens : c'est le *minimum de déviation*.

On repère la position de l'image en pointant avec la lunette (position L_1 de la figure 14.5).

On procède de la même manière en éclairant l'autre face avec le collimateur (position L_2 de la figure 14.5).

L'angle entre les deux positions de la lunette est égal à $2D_m$. On déduit alors n de l'équation 14.8.

On peut effectuer ces mesures avec une lampe spectrale contenant plusieurs longueurs d'onde comme la lampe à mercure ou au cadmium. Dans ce cas, on peut déterminer l'expression de n en fonction de la longueur d'onde et vérifier la loi de Cauchy (11.13).

3. Spectromètre à prisme

L'étude du spectre d'une lampe n'est pas très aisée avec un goniomètre car il faut relever la position de la lunette pour chaque raie pointée dans le spectre. On utilise alors plus couramment un appareil inspiré du goniomètre mais permettant de repérer les raies du spectre sur une graduation directement projetée sur le prisme et visible à travers la lunette : il s'agit d'un spectromètre à prisme (Cf. figure 14.6).

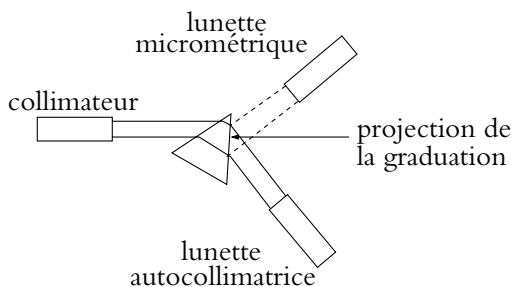


Figure 14.6 Disposition des instruments pour la spectroscopie à prisme.

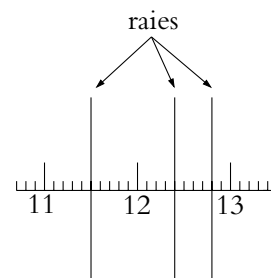


Figure 14.7 Superposition du spectre et de l'échelle micrométrique.

Il est constitué :

- d'un plateau sur lequel est posé un prisme ;
- d'un collimateur que l'on éclaire avec une lampe spectrale ;
- d'une lunette autocollimatrice ;
- d'une lunette micrométrique.

Il y a donc un instrument supplémentaire par rapport au goniomètre : la lunette micrométrique.

Son rôle est de projeter une échelle graduée qui se réfléchit sur une face du prisme et est vue superposée au spectre à travers la lunette autocollimatrice (Cf. figure 14.7). Une bague permet le réglage de la netteté. Il suffit alors de repérer les positions x de raies de longueurs d'onde λ connues, par exemple celles du mercure, sur la règle graduée. On trace une courbe d'étalonnage λ en fonction de x . On éclaire ensuite le collimateur avec une lampe dont on veut déterminer les longueurs d'onde. On repère ensuite les positions x des raies inconnues. On reporte ces positions sur la courbe d'étalonnage et on en déduit leurs longueurs d'onde.

DEUXIÈME PÉRIODE

Électrocinétique 2

Circuits linéaires en régime sinusoïdal forcé

15

Dans ce chapitre, on étudie le régime sinusoïdal forcé qui est un régime variable et permanent. On présente les signaux sinusoïdaux ainsi que l'outil essentiel pour l'étude des systèmes linéaires soumis à ce type de signaux : la notation complexe. On justifie son intérêt sur l'exemple d'un circuit simple : le circuit RC et on généralise les résultats énoncés dans le chapitre de la première période sur les régimes continus au cas du régime sinusoïdal notamment en définissant les notions d'impédance et d'admittance.

1. Caractéristiques d'un signal sinusoïdal

1.1 Définition

Un signal sinusoïdal est un signal dépendant du temps t (on dit qu'il est variable) et dont l'expression en fonction de t est :

$$x(t) = X \cos(\omega t + \varphi)$$

ou

$$x(t) = X \sin(\omega t + \varphi)$$

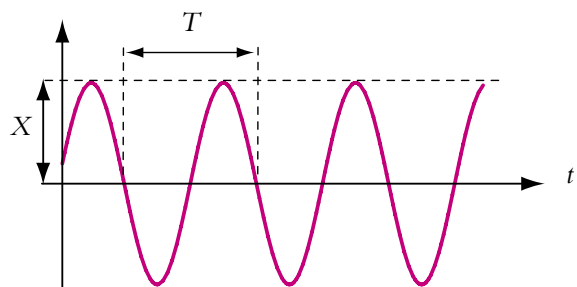


Figure 15.1 Signal sinusoïdal.

On ne précise pas la nature du signal : il peut s'agir d'une intensité, d'une tension ou de toute autre grandeur physique. Dans la suite, on utilisera la fonction cosinus ; tout ce qui sera dit sera transposable au cas d'une fonction sinus.

1.2 Amplitude

La quantité X est appelée *amplitude*.

C'est une constante, caractéristique d'un signal sinusoïdal.

1.3 Pulsation, fréquence et période

La grandeur ω est appelée *pulsation*, elle s'exprime en radians par seconde (rad.s^{-1}).

Il s'agit également d'une constante, autre caractéristique d'un signal sinusoïdal. Physiquement, cette grandeur fournit la notion de rapidité avec laquelle le signal reprend une valeur identique à celle qu'il avait à un instant antérieur avec les mêmes valeurs antérieures et postérieures. Cela correspond à l'idée de périodicité du signal.

De ce fait, il est possible de remplacer cette caractéristique par deux grandeurs également constantes qui s'expriment par des relations simples en fonction de ω :

- la *fréquence* f (en hertz) telle que :

$$f = \frac{\omega}{2\pi}$$

- la *période* T (en seconde) telle que :

$$T = \frac{2\pi}{\omega}$$

On retiendra aussi la relation entre la fréquence et la période :

$$f = \frac{1}{T}$$

Il sera parfaitement équivalent de donner la pulsation, la fréquence ou la période du signal. Ces trois grandeurs décrivent la même notion physique de périodicité.

1.4 Phases

On distingue deux types de phase :

- la *phase instantanée* qui est, par définition, la quantité $\omega t + \varphi$,
- la *phase initiale*, φ .

Dans la suite, on ne parlera que de la phase initiale sauf indication contraire et on omettra le terme « initial » à moins qu'il puisse y avoir un risque de confusion.

Cette notion de phase permet de fixer l'origine des temps.

1.5 Ensemble des caractéristiques d'un signal sinusoïdal

Un signal sinusoïdal $x(t) = X \cos(\omega t + \varphi)$ est complètement déterminé par son amplitude X , sa pulsation ω (ou sa période T ou sa fréquence f) et sa phase initiale φ . On donnera donc ces trois informations pour définir un tel signal. Les autres grandeurs se déduisent facilement de ces trois quantités.



Un signal sinusoïdal peut également être donné sous la forme :

$$x(t) = A \cos \omega t + B \sin \omega t$$

Dans ce cas, son amplitude est :

$$X = \sqrt{A^2 + B^2}$$

et sa phase initiale φ vérifie :

$$\cos \varphi = \frac{A}{\sqrt{A^2 + B^2}} \quad \text{et} \quad \sin \varphi = -\frac{B}{\sqrt{A^2 + B^2}}$$

Pour établir ces relations, il suffit d'écrire

$$x(t) = X \cos(\omega t + \varphi) = X \cos \varphi \cos \omega t - X \sin \varphi \sin \omega t$$

et d'identifier les termes en $\cos \omega t$ et en $\sin \omega t$ dans les deux expressions de $x(t)$.

1.6 Valeur efficace

On a également l'habitude de définir la valeur efficace¹ X_{eff} qui vaut dans le cas d'une grandeur sinusoïdale centrée (c'est-à-dire de valeur moyenne nulle) :

$$X_{\text{eff}} = \frac{X}{\sqrt{2}}$$

On reviendra dans les chapitres sur la puissance et l'instrumentation sur cette notion.

1.7 Différence de phase entre deux signaux synchrones

On considère deux signaux de même pulsation (ils sont dits *synchrones*) :

$$x(t) = X \cos(\omega t + \varphi) \quad \text{et} \quad y(t) = Y \cos(\omega t + \psi)$$

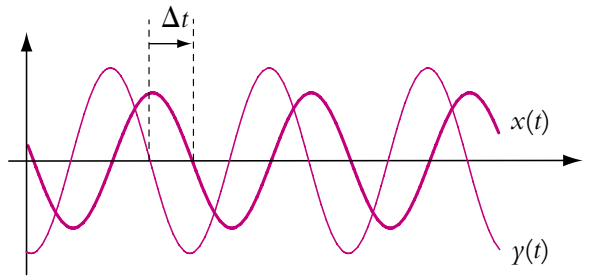


Figure 15.2 Différence de phase entre deux signaux synchrones.

¹De manière générale, la valeur efficace X_{eff} d'un signal $x(t)$ de période T est définie par :

$$X_{\text{eff}} = \sqrt{\frac{1}{T} \int_{t_0}^{t_0+T} (x(t))^2 dt}$$

Le signal $y(t)$ passe par son maximum (par exemple) avant le signal $x(t)$: il est *en avance* de phase sur $x(t)$. Le déphasage $\Delta\varphi$ de $y(t)$ par rapport $x(t)$ est $\Delta\varphi = \psi - \varphi = \omega\Delta t$.

► **Remarque :** Cette définition suppose que le déphasage $\Delta\varphi$ soit compris entre $-\pi$ et $+\pi$. On pourrait, en effet, considérer que c'est $x(t)$ qui est en avance de phase de $2\pi - \Delta\varphi$ sur $y(t)$.

2. Intérêt de la notation complexe

2.1 Exemple du circuit R, C

On considère le circuit suivant :

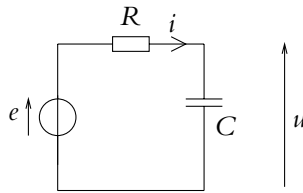


Figure 15.3 Circuit RC.

où e est une tension sinusoïdale $e = E_0 \cos \omega t$. On cherche à déterminer l'expression de l'intensité $i(t)$ et de la tension $u(t)$.

La loi des mailles donne :

$$e = Ri + u$$

Or la relation entre l'intensité du courant traversant un condensateur et la tension à ses bornes s'écrit :

$$i = C \frac{du}{dt}$$

et on obtient l'équation différentielle suivante :

$$\frac{du}{dt} + \frac{1}{RC}u = \frac{1}{RC}e$$

soit

$$\frac{du}{dt} + \frac{u}{\tau} = \frac{e}{\tau}$$

en posant $\tau = RC$.

Il s'agit d'une équation différentielle linéaire du premier ordre à coefficients constants dont la solution s'écrit comme la somme de :

- la solution générale² $u_H(t)$ de l'équation homogène associée $\frac{du}{dt} + \frac{u}{\tau} = 0$:

$$u_H(t) = U_0 \exp\left(-\frac{t}{\tau}\right)$$

- une solution particulière $u_P(t)$ qu'on cherche sous une forme analogue à celle du second membre :

$$u_P(t) = A \cos \omega t + B \sin \omega t$$

On doit donc déterminer les valeurs de A et B : on dérive l'expression $u_P(t)$ et on reporte dans l'équation différentielle :

$$\left(-A\omega + \frac{B}{\tau}\right) \sin \omega t + \left(\frac{A}{\tau} + B\omega\right) \cos \omega t = \frac{E_0}{\tau} \cos \omega t$$

En identifiant les coefficients du cosinus et du sinus, on obtient :

$$\begin{cases} \omega B + \frac{A}{\tau} = \frac{E_0}{\tau} \\ -\omega A + \frac{B}{\tau} = 0 \end{cases}$$

On en déduit :

$$A = \frac{E_0}{1 + (\tau\omega)^2} \quad \text{et} \quad B = \frac{\tau\omega E_0}{1 + (\tau\omega)^2}$$

La solution particulière cherchée s'écrit alors :

$$u_P(t) = \frac{E_0}{1 + (\tau\omega)^2} (\cos \omega t + \tau\omega \sin \omega t)$$

La solution générale est donc :

$$u(t) = u_H(t) + u_P(t) = U_0 \exp\left(-\frac{t}{\tau}\right) + \frac{E_0}{1 + (\tau\omega)^2} (\cos \omega t + \tau\omega \sin \omega t)$$

Au bout d'un temps égal à quelques τ , on a : $|u_H(t)| \ll |u_P(t)|$ et on peut écrire :

$$u(t) \simeq u_P(t)$$

Il s'agit du régime sinusoïdal forcé ($u_H(t)$ correspond au régime libre qui a été étudié au chapitre sur les circuits soumis à un échelon de tension).

² L'analyse de cette solution a été détaillée dans le chapitre sur les circuits linéaires soumis à un échelon de tension dans le cours de première période.

► **Remarque :** $u_P(t)$ peut également s'écrire sous la forme :

$$u_P(t) = U_0 \cos(\omega t + \varphi)$$

avec

$$U_0 = \frac{E_0}{\sqrt{1 + (\tau\omega)^2}} \quad \text{et} \quad \begin{cases} \cos \varphi = \frac{1}{\sqrt{1 + (\tau\omega)^2}} > 0 \\ \sin \varphi = -\frac{\tau\omega}{\sqrt{1 + (\tau\omega)^2}} < 0 \end{cases}$$

soit

$$U_0 = \frac{E_0}{\sqrt{1 + (\tau\omega)^2}} \quad \text{et} \quad \begin{cases} \varphi \in \left[-\frac{\pi}{2}, 0\right] \\ \tan \varphi = -\tau\omega \end{cases}$$



Attention, la seule donnée de $\tan \varphi$ ne suffit pas à déterminer φ à 2π près (la fonction tangente est de période π), il faut donc une donnée supplémentaire permettant de déterminer dans quel intervalle de longueur π se trouve φ : le signe du cosinus et/ou du sinus permet de lever l'indétermination sur φ .

2.2 Méthode fastidieuse et risquée

L'obtention de cette solution nécessite un nombre important de calculs qui, tout en étant parfaitement réalisables, peuvent conduire à des erreurs. On a pourtant considéré ici un circuit simple ne comprenant que deux dipôles et déjà les calculs sont assez longs. On imagine aisément la lourdeur de la résolution pour un système plus complexe.

L'objet de ce chapitre est de mettre en place une approche plus facile à utiliser et limitant les risques d'erreur de manière à ne pas devoir recourir à la méthode qui vient d'être exposée. Il s'agit d'utiliser la notation complexe qu'on va présenter dans le paragraphe suivant.

3. Notation complexe

3.1 Signaux complexes ou analytiques associés

À la grandeur sinusoïdale $x(t) = X \cos(\omega t + \varphi)$, on peut associer un signal complexe $\underline{x}(t)$ défini par :

$$\underline{x}(t) = X \exp(j(\omega t + \varphi))$$

On souligne le symbole x pour préciser qu'il s'agit d'une notation complexe et non réelle.

On notera qu'en électricité, on préfère habituellement désigner par j la racine carrée de -1 au lieu de i comme c'est le cas en mathématiques. j correspond donc au nombre

complexe de module 1 et d'argument $\frac{\pi}{2}$ vérifiant $j^2 = -1$. Cette notation permet d'éviter toute confusion avec l'intensité i .

On peut écrire :

$$\underline{x}(t) = (X \exp(j\varphi)) \exp(j\omega t)$$

On appelle *amplitude complexe* la grandeur complexe :

$$\underline{X} = X \exp(j\varphi)$$

qui ne dépend pas du temps. Elle contient toutes les informations sur l'amplitude et la phase initiale du signal.

À partir du signal complexe, on revient facilement au signal réel de départ :

$$x(t) = \operatorname{Re}(\underline{x}(t))$$

et les caractéristiques du signal sont obtenues :

- en prenant le module de $\underline{X}$ ou de $\underline{x}(t)$ pour l'amplitude :

$$X = |\underline{x}(t)| = |\underline{X}|$$

- en prenant l'argument de $\underline{X}$ pour la phase initiale :

$$\varphi = \operatorname{Arg}(\underline{X})$$

- en prenant l'argument de $\underline{x}(t)$ pour la phase instantanée :

$$\omega t + \varphi = \operatorname{Arg}(\underline{x}(t))$$

► **Remarque** : On se place ici dans le cas d'équations différentielles linéaires à second membre sinusoïdal du type $E_0 \cos \omega t$ dont on cherche une solution particulière sous la forme $u_{p1}(t) = U_0 \cos(\omega t + \varphi)$.

Quand on translate l'origine des temps d'un quart de période ($t \rightarrow t - \frac{T}{4}$) dans l'équation différentielle initiale (E_1), seul le second membre change. Il devient $E_0 \sin \omega t$ et une solution particulière de cette nouvelle équation différentielle (E_2) est $u_{p2}(t) = U_0 \sin(\omega t + \varphi)$.

La somme des équations différentielles ($\underline{E} = \langle E_1 \rangle + j \langle E_2 \rangle$) a pour second membre $E_0 \exp(j\omega t)$ et une solution particulière de ($\underline{E}$) est $\underline{u}_p = u_{p1}(t) + ju_{p2}(t)$. Cette propriété, liée à la linéarité des équations étudiées, justifie l'utilisation de la notation complexe. On cherchera donc $\underline{x}(t)$, solution de ($\underline{E}$).

Pour les règles de calcul sur les modules et les arguments d'un produit ou d'un quotient de nombres complexes, le lecteur est invité à se reporter à son cours de mathématiques de première période.

3.2 Représentation de Fresnel

À cette notation complexe, on associe également une représentation vectorielle dite *représentation de Fresnel*. Il s'agit de tracer dans le plan complexe le vecteur d'origine O le centre du repère et correspondant au complexe $\underline{x}(t)$, c'est-à-dire le vecteur $\overrightarrow{OM}$ d'affixe $\underline{x}(t)$ tel que :

$$\begin{cases} OM = X \\ (\widehat{Ox, OM}) = \omega t + \varphi \end{cases}$$

Cela revient à associer à la grandeur sinusoïdale $x(t) = X \cos(\omega t + \varphi)$ un vecteur tournant dans le sens trigonométrique à la vitesse angulaire constante ω autour de l'origine et dont le module vaut l'amplitude X . Ce vecteur est appelé *vecteur de Fresnel*.

En général, on le représente à l'instant $t = 0$.

Une telle représentation offre l'avantage de donner une visualisation géométrique. Cela pourra permettre, à condition que le schéma soit clair, d'alléger les calculs en « lisant » la figure.

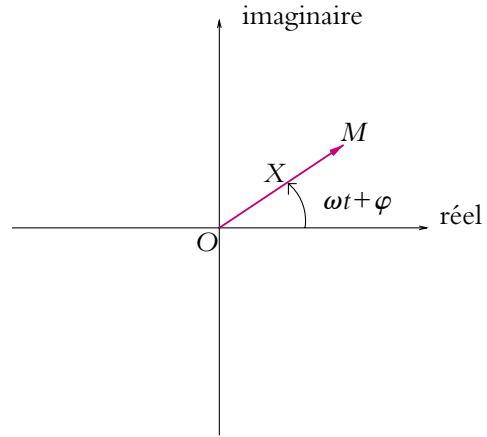


Figure 15.4 Représentation de Fresnel d'un signal sinusoïdal.

3.3 Addition de signaux de même pulsation

Soient deux grandeurs sinusoïdales :

$$x_1(t) = X_1 \cos(\omega t + \varphi_1) \quad \text{et} \quad x_2(t) = X_2 \cos(\omega t + \varphi_2)$$

de même pulsation ω . Les grandeurs complexes associées sont respectivement $\underline{x}_1(t) = X_1 \exp(j\varphi_1) \exp(j\omega t)$ et $\underline{x}_2(t) = X_2 \exp(j\varphi_2) \exp(j\omega t)$ et on note $\overrightarrow{OM_1}$ et $\overrightarrow{OM_2}$ les vecteurs de Fresnel associés.

a) En notation complexe

Prendre la partie réelle d'une somme de nombres complexes revient à additionner les parties réelles de ces nombres complexes, on en déduit que

$$\operatorname{Re}(\underline{x}_1(t) + \underline{x}_2(t)) = \operatorname{Re}(\underline{x}_1(t)) + \operatorname{Re}(\underline{x}_2(t)) = x_1(t) + x_2(t)$$

La grandeur $X_1 \exp(j\varphi_1) \exp(j\omega t) + X_2 \exp(j\varphi_2) \exp(j\omega t)$ est donc le nombre complexe associé à la somme des signaux réels $x_1(t)$ et $x_2(t)$. Comme les deux signaux ont

même pulsation ω , le terme $\exp(j\omega t)$ se met en facteur et $\underline{x}(t) = \underline{x}_1(t) + \underline{x}_2(t)$ a pour amplitude $\underline{X}_1 + \underline{X}_2 = X_1 \exp(j\varphi_1) + X_2 \exp(j\varphi_2)$.

b) En représentation de Fresnel

On peut associer à ce nombre complexe une représentation de Fresnel $\overrightarrow{OM}$ qui est la somme vectorielle de $\overrightarrow{OM}_1$ et de $\overrightarrow{OM}_2$. En effet, dans le plan complexe, la somme des vecteurs représentant deux nombres complexes est un vecteur représentant la somme des deux nombres complexes.

On peut noter que cela nécessite l'indéformabilité du triangle OM_1M_2 au cours du temps. Si ce n'est pas le cas, la somme vectorielle ne pourra donner un vecteur de module constant et il ne sera pas possible de définir l'amplitude complexe. Or ce triangle ne sera indéformable que si les deux vecteurs $\overrightarrow{OM}_1$ et $\overrightarrow{OM}_2$ tournent autour de O à la même vitesse angulaire. Cela impose aux deux signaux d'avoir même pulsation, ce qui a été supposé au début du paragraphe.

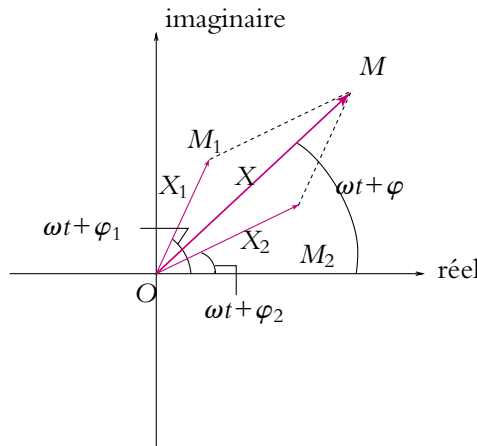


Figure 15.5 Addition de nombres complexes en représentation de Fresnel.

3.4 Dérivation de signaux

On cherche à représenter la dérivée de $x(t) = X \cos(\omega t + \varphi)$ qui vaut :

$$\frac{dx}{dt} = -\omega X \sin(\omega t + \varphi) = \omega X \cos\left(\omega t + \varphi + \frac{\pi}{2}\right)$$

par la grandeur complexe et son vecteur de Fresnel associés.

La grandeur complexe associée est :

$$X\omega \exp\left(j\left(\omega t + \varphi + \frac{\pi}{2}\right)\right) = \exp\left(j\frac{\pi}{2}\right) \omega X \exp(j(\omega t + \varphi)) = j\omega \underline{x}(t)$$

La dérivation d'un signal $x(t)$ se traduit par la multiplication par $j\omega$ pour la représentation complexe associée.

Le vecteur de Fresnel a pour module ωX et pour phase initiale $\varphi + \frac{\pi}{2}$. Il s'agit donc du vecteur représentant $x(t)$ auquel on a fait subir une homothétie de centre O et de rapport ω puis une rotation de centre O et d'angle $\frac{\pi}{2}$ dans le sens trigonométrique.

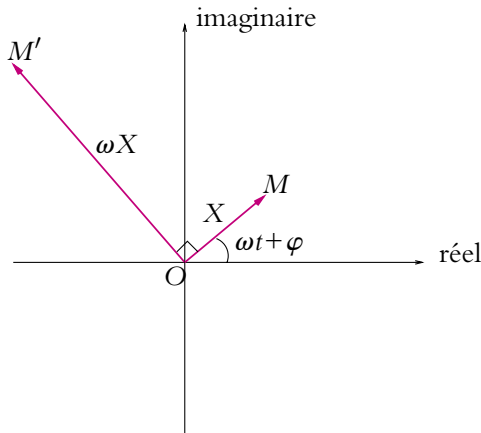


Figure 15.6 Dérivation d'un nombre complexe en représentation de Fresnel.

3.5 Intégration de signaux

On cherche à représenter le signal obtenu par intégration de $x(t) = X \cos(\omega t + \varphi)$ qui vaut :

$$\int x(t)dt = \frac{X}{\omega} \sin(\omega t + \varphi) = \frac{X}{\omega} \cos\left(\omega t + \varphi - \frac{\pi}{2}\right)$$

par la grandeur complexe et son vecteur de Fresnel associés.

La grandeur complexe associée est :

$$\frac{X}{\omega} \exp\left(j\left(\omega t + \varphi - \frac{\pi}{2}\right)\right) = \exp\left(-j\frac{\pi}{2}\right) \frac{1}{\omega} X \exp(j(\omega t + \varphi)) = \frac{1}{j\omega} \underline{x(t)}$$

L'intégration d'un signal $x(t)$ se traduit par la multiplication par $\frac{1}{j\omega}$ pour la représentation complexe associée.

Le vecteur de Fresnel a pour module $\frac{X}{\omega}$ et pour phase initiale $\varphi - \frac{\pi}{2}$. Il s'agit donc du vecteur représentant $x(t)$ auquel on a fait subir une homothétie de centre O et de rapport $\frac{1}{\omega}$ puis une rotation de centre O et d'angle $-\frac{\pi}{2}$ dans le sens trigonométrique (ou $\frac{\pi}{2}$ dans le sens des aiguilles d'une montre).

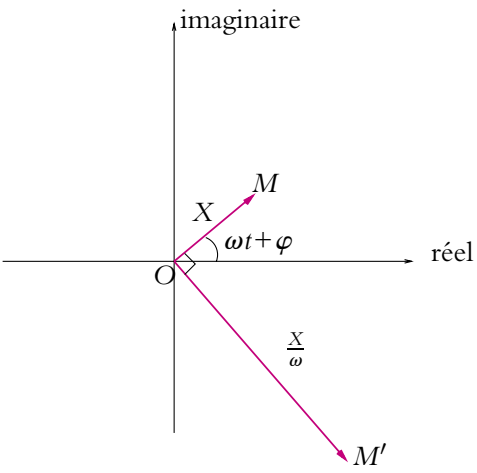


Figure 15.7 Intégration d’un nombre complexe en représentation de Fresnel.

On en déduit l’analogie suivante entre l’espace réel et l’espace complexe :

	signaux réels	représentation complexe
dérivation	$\frac{d}{dt}$	$\times j\omega$
intégration	$\int dt$	$\times \frac{1}{j\omega}$

3.6 Utilisation de la notation complexe pour l’étude du circuit R, C

La loi des mailles donne :

$$Ri(t) + u(t) = E_0 \cos \omega t$$

soit en passant en notation complexe :

$$R\underline{i}(t) + \underline{u}(t) = E_0 \exp(j\omega t)$$

La relation $i = C \frac{du}{dt}$ s’écrit en notation complexe : $\underline{i}(t) = jC\omega \underline{u}(t)$. On en déduit la relation :

$$\underline{u}(t) (1 + jRC\omega) = \underline{u}(t) (1 + j\tau\omega) = E_0 \exp(j\omega t) \quad \text{soit} \quad \underline{u}(t) = \frac{E_0 \exp(j\omega t)}{1 + j\tau\omega}$$

Le module de $\underline{u}(t)$ vaut :

$$U = \frac{|E_0 \exp(j\omega t)|}{|1 + j\tau\omega|} = \frac{E_0}{\sqrt{1 + \tau^2\omega^2}}$$

et le déphasage par rapport à $e(t)$:

$$\varphi = \text{Arg}(E_0) - \text{Arg}(1 + j\tau\omega) = -\text{Arg}(1 + j\tau\omega)$$

soit

$$\begin{cases} \tan \varphi = -\frac{\operatorname{Im}(1 + j\tau\omega)}{\operatorname{Re}(1 + j\tau\omega)} = -\tau\omega \\ \varphi \in \left[-\frac{\pi}{2}, 0\right] \end{cases}$$

On rappelle que la tangente d'un angle ne suffit pas pour connaître sa valeur à 2π près : on ne l'a qu'à π près. Il faut donc préciser l'intervalle de longueur π dans lequel se trouve cet angle, ce qui se détermine à partir du signe du sinus et/ou du cosinus. Ici on a $\sin \varphi < 0$ et $\cos \varphi > 0$ donc $\varphi \in \left[-\frac{\pi}{2}, 0\right]$. On peut donc écrire ici : $\varphi = -\arctan \tau\omega$. On obtient finalement l'expression :

$$\underline{u}(t) = \frac{E_0}{\sqrt{1 + \tau^2\omega^2}} \exp(j(\omega t - \arctan \tau\omega))$$

et en notation réelle :

$$u(t) = \frac{E_0}{\sqrt{1 + \tau^2\omega^2}} \cos(\omega t - \arctan \tau\omega)$$

On obtient beaucoup plus rapidement et beaucoup plus aisément la solution déterminée au début de ce chapitre. Cela montre l'intérêt de l'utilisation de la notation complexe pour l'étude des circuits en régime sinusoïdal forcé.

4. Lois de Kirchhoff en notation complexe

4.1 Validité des lois du continu pour le régime sinusoïdal

La première chose à faire est de justifier que les lois et les théorèmes vus lors de l'étude du régime continu restent valables en régime sinusoïdal.

Cette justification repose entièrement sur l'approximation des régimes quasi-stationnaires et sur la linéarité des équations comme cela a déjà été souligné précédemment. Dans le cadre de cette approximation, on néglige les phénomènes de propagation pour l'intensité et la tension. Cela revient à dire qu'une modification de l'intensité à une extrémité d'un fil de connexion se transmet instantanément à l'autre extrémité.

On peut donc énoncer l'ensemble des lois et des théorèmes déjà vus dans le cadre du régime continu pour le régime sinusoïdal. Il suffit de remplacer les grandeurs par les grandeurs instantanées correspondantes.

4.2 Loi des nœuds

Soit un nœud N , point de concours de n courants $i_1(t)$, $i_2(t)$, ... $i_n(t)$. On accepte que les sens conventionnels des intensités puissent partir ou arriver du nœud : on pose $\epsilon_k = \pm 1$ en fonction du sens conventionnel choisi pour la branche k .

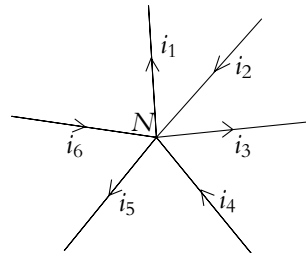


Figure 15.8 Loi des nœuds.

En régime continu, on écrivait :

$$\sum_{k=1}^n \epsilon_k i_k = 0$$

qui se traduit en valeurs instantanées dans le cas des régimes variables par :

$$\sum_{k=1}^n \epsilon_k i_k(t) = 0$$

Or $i_k(t)$ en régime sinusoïdal peut s'obtenir à partir de la grandeur complexe associée : $i_k(t) = \text{Re}(\underline{i_k}(t))$.

On en déduit : $\sum_{k=1}^n \text{Re}(\epsilon_k \underline{i_k}(t)) = 0$.

D'après la remarque du paragraphe 3.1, les grandeurs complexes vérifient la même équation, à savoir :

$$\sum_{k=1}^n \epsilon_k \underline{i_k}(t) = 0$$

Les courants ayant tous même pulsation, on peut simplifier cette équation par $\exp(j\omega t)$. On obtient :

$$\sum_{k=1}^n \epsilon_k \underline{I_k} = 0$$

Il s'agit de la loi des nœuds traduite en notation complexe pour les signaux sinusoïdaux.

4.3 Loi des mailles

Soit une maille constituée de n branches.

En régime continu, on écrivait :

$$\sum_{k=1}^n u_k = 0$$

qui se traduit en valeurs instantanées dans le cas des régimes variables par :

$$\sum_{k=1}^n u_k(t) = 0$$

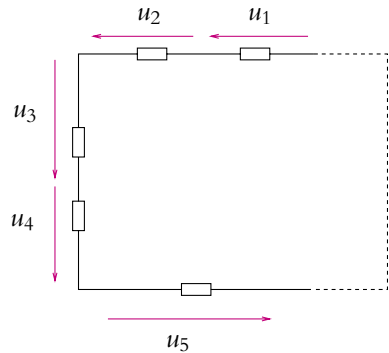


Figure 15.9 Loi des mailles.

Or $u_k(t)$ en régime sinusoïdal peut s'obtenir à partir de la grandeur complexe associée : $u_k(t) = \text{Re}(\underline{u}_k(t))$.

On en déduit : $\sum_{k=1}^n \text{Re}(\underline{u}_k(t)) = 0$.

Pour les mêmes raisons qu'avec la loi des nœuds, on peut écrire à chaque instant :

$$\sum_{k=1}^n \underline{u}_k(t) = 0$$

Les tensions ayant tous même pulsation, on peut simplifier cette équation par $\exp(j\omega t)$. On obtient :

$$\sum_{k=1}^n \underline{U}_k = 0$$

Il s'agit de la loi des mailles traduite en notation complexe pour les signaux sinusoïdaux.

4.4 Lois de Kirchhoff

Les lois de Kirchhoff ont donc la même formulation en notation complexe pour les signaux sinusoïdaux qu'en notation réelle pour le régime continu. En effet, la méthode des mailles (respectivement la méthode des nœuds) qui permet de déterminer l'intensité dans toutes les branches (respectivement la tension entre tous les couples de nœuds) repose uniquement sur l'expression de lois des nœuds et/ou de lois des mailles.

Il suffit donc d'utiliser la notation complexe pour avoir la résolution correspondante dans le cas sinusoïdal : on détermine les amplitudes complexes des intensités dans les différentes branches (respectivement les amplitudes complexes des tensions aux bornes des différents dipôles) d'un réseau de dipôles linéaires à partir de la méthode des mailles (respectivement de la méthode des nœuds).

5. Impédance et admittance complexes

5.1 Modèles de Thévenin et de Norton

Au chapitre 2, on a défini de manière générale un dipôle linéaire comme un dipôle dont l'intensité $i(t)$ qui le traverse et la tension $u(t)$ à ses bornes sont reliées par une relation différentielle linéaire à coefficients constants :

$$\sum_{k=0}^K a_k \frac{d^k u}{dt^k} + \sum_{l=1}^L b_l \frac{d^l i}{dt^l} = f(t)$$

Ici on est en régime sinusoïdal donc $f(t) = F \cos(\omega t + \psi)$.

Compte tenu de la linéarité des équations différentielles, il est parfaitement justifié d'adopter la notation complexe.

La dérivation d'une grandeur complexe revenant à la multiplier par $j\omega$, la relation précédente peut donc s'écrire sous la forme :

$$\sum_{k=0}^K a_k (j\omega)^k \underline{u}(t) + \sum_{l=1}^L b_l (j\omega)^l \underline{i}(t) = \underline{f}(t)$$

soit :

$$\underline{u}(t) = \frac{\underline{f}(t)}{\sum_{k=0}^K a_k (j\omega)^k} - \frac{\sum_{l=1}^L b_l (j\omega)^l}{\sum_{k=0}^K a_k (j\omega)^k} \underline{i}(t)$$

qu'on peut noter :

$$\underline{u}(t) = \underline{E}(t) - \underline{Z} \underline{i}(t)$$

Cette expression est analogue à $u = E - Ri$ qu'on a pu utiliser en régime continu.

On obtient ainsi le modèle de Thévenin du dipôle linéaire en notation complexe.

Par la même transformation que celle utilisée en régime continu, on obtient le modèle de Norton du même dipôle linéaire :

$$\underline{i}(t) = \underline{i}_0(t) - \underline{Y} \underline{u}(t) \quad \text{avec} \quad \underline{Y} = \frac{1}{\underline{Z}} \quad \text{et} \quad \underline{i}_0(t) = \frac{\underline{E}(t)}{\underline{Z}}$$

5.2 Définition de l'impédance et de l'admittance complexe pour un dipôle linéaire

a) Impédance complexe

On définit $\underline{Z}$ comme l'*impédance complexe* du dipôle considéré.

Dans le cas des modèles de Thévenin et de Norton précédents, il s'agit de l'impédance interne des générateurs.

Pour les autres dipôles linéaires pour lesquels $\underline{E}(t) = 0$, on obtient une « relation de proportionnalité » en notation complexe :

$$\underline{U} \exp(j\omega t) = \underline{Z} \underline{I} \exp(j\omega t)$$

soit

$$\underline{U} = \underline{Z} \underline{I} \quad (15.1)$$

analogue à $u = Ri$ pour une résistance.

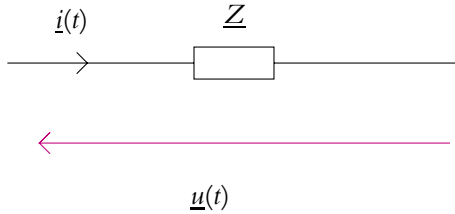


Figure 15.10 Représentation d'une impédance.

Ce type de relation avait disparu dans la généralisation de la définition d'un dipôle linéaire à l'aide d'une équation différentielle. La notation complexe permet de retrouver la même forme en introduisant l'équivalent complexe de la résistance : l'impédance.

► **Remarque :** Cette relation est également valable pour les grandeurs complexes associées $\underline{u}(t) = U \exp(j\varphi_u) \exp(j\omega t)$ et $\underline{i}(t) = I \exp(j\varphi_i) \exp(j\omega t)$:

$$\underline{u}(t) = \underline{Z} \underline{i}(t) \quad (15.2)$$

Ces deux relations (15.1) et (15.2) portent le nom de loi d'Ohm complexe.

b) Admittance complexe

De la même façon que l'on a défini la conductance comme l'inverse d'une résistance, on appelle *admittance complexe* l'inverse de l'impédance complexe. Habituellement on la note :

$$\underline{Y} = \frac{1}{\underline{Z}}$$

5.3 Impédance et déphasage

L'impédance complexe peut s'écrire comme tout complexe sous la forme :

$$\underline{Z} = Z \exp(j\psi)$$

On appelle *impédance* la grandeur Z (réelle) correspondant au module de l'impédance complexe $\underline{Z}$. Il faut faire particulièrement attention aux quantités qu'on manipule et savoir à tout moment s'il s'agit d'une grandeur complexe ou d'une grandeur réelle.

On peut noter que :

$$\underline{Z} = \frac{U}{I} = \frac{U \exp(j\varphi_u)}{I \exp(j\varphi_i)} = \frac{U}{I} \exp(j(\varphi_u - \varphi_i))$$

donc l'impédance Z est le rapport des amplitudes de la tension et de l'intensité :

$$Z = \frac{U}{I}$$

De même, l'argument ψ de $\underline{Z}$ est égal au déphasage de la tension u par rapport à l'intensité i :

$$\psi = \varphi_u - \varphi_i$$

On note que la connaissance de l'intensité (respectivement de la tension) complexe et de l'impédance complexe permet de déterminer complètement la tension (respectivement de l'intensité) complexe.

5.4 Résistance et réactance

On peut également utiliser l'écriture d'un nombre complexe sous forme partie réelle - partie imaginaire :

$$\underline{Z} = R + jX = Z \cos \psi + jZ \sin \psi$$

On appelle :

- *résistance* la grandeur :

$$R = Z \cos \psi$$

- *réactance* la grandeur :

$$X = Z \sin \psi$$

5.5 Exemples d'impédances

Dans tout ce paragraphe, on utilise la convention récepteur. Si on a besoin de la convention générateur, il suffit d'ajouter un signe $-$ aux relations qui vont être données.

a) Résistor de résistance R

La relation temporelle instantanée $u(t) = Ri(t)$ s'écrit en notation complexe : $\underline{u}(t) = R\underline{i}(t)$, ou encore : $\underline{U} = R\underline{I}$.

Par identification, on en déduit que l'impédance complexe d'une résistance est réelle et vaut :

$$\underline{Z} = R$$

On remarque qu'une résistance n'introduit pas de déphasage entre $u(t)$ et $i(t)$.

b) Bobine d'inductance L

On considère une bobine dont on néglige la résistance interne, on ne s'intéresse qu'à l'inductance dont la relation instantanée entre intensité et tension s'exprime par : $u(t) = L \frac{di}{dt}$ ce qui se traduit en notation complexe par : $\underline{u}(t) = jL\omega \underline{i}(t)$, ou encore par : $\underline{U} = jL\omega \underline{I}$. On en déduit l'impédance complexe :

$$\underline{Z} = jL\omega$$

Il s'agit d'une impédance imaginaire pure : on parle de réactance pure. L'argument vaut $\frac{\pi}{2}$ et l'intensité et la tension seront en quadrature de phase, la tension étant en avance sur l'intensité.

En basse fréquence, l'impédance d'une bobine tend vers 0 : la bobine est donc équivalente à un fil.

En haute fréquence, l'impédance d'une bobine tend vers $+\infty$: la bobine est donc équivalente à un interrupteur ouvert.

c) Condensateur de capacité C

On considère un condensateur de capacité C . La relation $i(t) = C \frac{du}{dt}$ se traduit en notation complexe par : $\underline{i}(t) = jC\omega \underline{u}(t)$ soit $\underline{u}(t) = \frac{1}{jC\omega} \underline{i}(t)$ ou encore par $\underline{U} = \frac{1}{jC\omega} \underline{I}$. On en déduit l'impédance associée

$$\underline{Z} = \frac{1}{jC\omega}$$

L'impédance est également imaginaire pure avec un argument de $-\frac{\pi}{2}$. L'intensité est en quadrature avance sur la tension. Il s'agit encore d'une réactance pure.

En basse fréquence, l'impédance d'un condensateur tend vers $+\infty$: le condensateur est équivalent à un interrupteur ouvert.

En haute fréquence, l'impédance d'un condensateur tend vers 0 : le condensateur est équivalent à un fil.

Pour ces trois dipôles, on retrouve bien une relation linéaire en complexe, ce qui justifie, si besoin était, leur classification dans les dipôles linéaires.

5.6 Association de dipôles linéaires

On a vu que les lois de Kirchhoff étaient valables en notation complexe. On va donc déterminer les relations d'association de dipôles en notation complexe par analogie avec l'association de résistances en régime continu.

a) Association en série

Soient n impédances notées $\underline{Z}_k$ placées en série.

L'association en série d'impédances est une impédance qui s'obtient comme la somme des impédances :

$$\underline{Z} = \underline{Z}_1 + \underline{Z}_2 + \dots + \underline{Z}_n$$

autrement dit, associées en série, les impédances s'ajoutent.

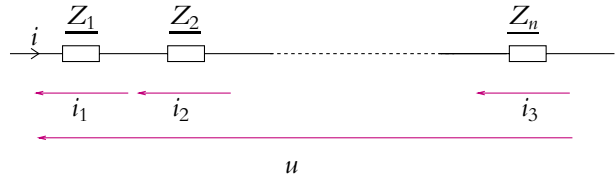


Figure 15.11 Association en série d'impédances.

b) Association en parallèle

Soient n impédances notées $\underline{Z}_k$ placées en parallèle.

L'association en parallèle d'impédances est une impédance dont la valeur s'obtient comme l'inverse de la somme de l'inverse des impédances :

$$\frac{1}{\underline{Z}} = \frac{1}{\underline{Z}_1} + \frac{1}{\underline{Z}_2} + \dots + \frac{1}{\underline{Z}_n}$$

ou encore

$$\underline{Y} = \underline{Y}_1 + \underline{Y}_2 + \dots + \underline{Y}_n$$

autrement dit, associées en parallèle, ce sont les admittances qui s'ajoutent.

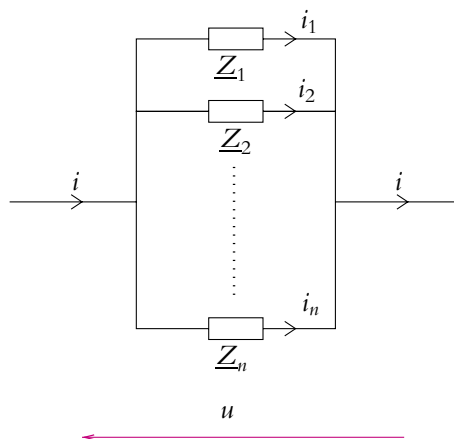


Figure 15.12 Association en parallèle d'impédances.

6. Loi des nœuds exprimée en termes de potentiels

6.1 Expression de la loi des nœuds en termes de potentiels

Soit un nœud N où convergent plusieurs branches indicées par k , constituées de dipôles linéaires passifs d'impédance complexe $\underline{Z}_k$ ou d'admittance complexe $\underline{Y}_k = \frac{1}{\underline{Z}_k}$.

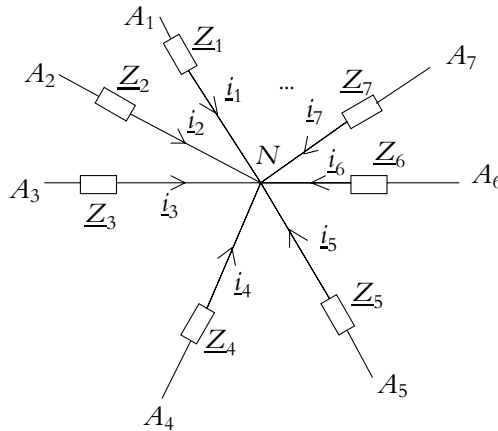


Figure 15.13 Loi des nœuds en termes de potentiels.

La loi des nœuds en N s'écrit :

$$\sum_k i_k = 0$$

où i_k est l'intensité complexe circulant dans la k -ième branche orientée conventionnellement vers le nœud N .

Or

$$i_k = \frac{V_k - V_N}{\underline{Z}_k}$$

en convention récepteur et en notant V_k le potentiel complexe du point A_k et V_N celui du nœud. La loi des nœuds devient :

$$\sum_k \frac{V_k - V_N}{\underline{Z}_k} = 0$$

Il s'agit de la loi des nœuds exprimée en termes de potentiels.

6.2 Complément : théorème de Millman

On peut en déduire une autre formulation qui pourra être utile notamment dans les montages à amplificateurs opérationnels qui seront vus ultérieurement.

a) Démonstration du théorème de Millman

À partir de la relation établie au paragraphe précédent, on peut écrire :

$$\sum_k \frac{V_k}{Z_k} - V_N \sum_k \frac{1}{Z_k} = 0$$

puis :

$$V_N = \frac{\sum_k \frac{V_k}{Z_k}}{\sum_k \frac{1}{Z_k}}$$

ou en utilisant les admittances complexes :

$$V_N = \frac{\sum_k Y_k V_k}{\sum_k Y_k}$$

Ce résultat porte le nom de théorème de Millman qui est simplement une reformulation de la loi des nœuds en termes de potentiels.

b) Généralisation du théorème de Millman

À partir de ce qui vient d'être fait, il est possible de proposer sans difficulté une généralisation de ce théorème au cas où on aurait des sources idéales de courant complexe connectées au nœud N . Il faut alors tenir compte des intensités complexes liées à ces sources idéales de courant dans l'écriture de la loi des nœuds.

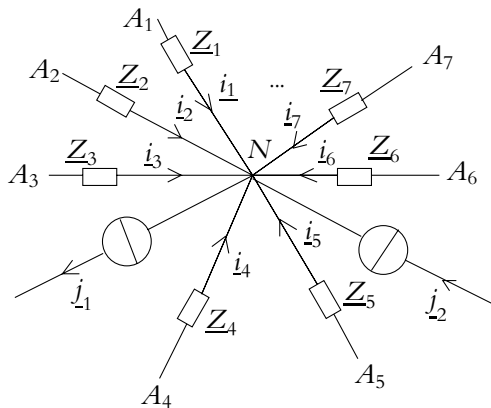


Figure 15.14 Théorème de Millman généralisé.

On distingue les branches où on a un générateur idéal de courant j_k de celles où on a une impédance Z_l . Dans la suite, les intensités j_k sont considérées comme algébriques : elles peuvent arriver ou partir du nœud. Il faudra en tenir compte en mettant en signe – devant les intensités quand le courant part du nœud.

La loi des nœuds s'écrit alors :

$$\sum_l i_l + \sum_k j_k = 0$$

Or comme précédemment :

$$\sum_l i_l = \sum_l \frac{V_l - V_N}{Z_l}$$

et :

$$\sum_l i_l + \sum_k j_k = \sum_l \frac{V_l}{Z_l} - V_N \sum_l \frac{1}{Z_l} + \sum_k j_k = 0$$

soit :

$$V_N = \frac{\sum_l \frac{V_l}{Z_l} + \sum_k j_k}{\sum_l \frac{1}{Z_l}} = \frac{\sum_l Y_l V_l + \sum_k j_k}{\sum_l Y_l}$$

Il s'agit donc du théorème de Millman généralisé où on prend en compte TOUTES les intensités complexes qu'elles soient issues d'une source idéale de courant ou qu'elles correspondent à un courant dans une impédance.



La conséquence immédiate de cette remarque est l'impossibilité d'utiliser ce théorème si l'intensité complexe du courant dans une seule branche est inconnue, sauf s'il s'agit de déterminer cette inconnue. Cela pourra poser un problème pour les sources idéales de tension dont on ne sait pas quelle intensité les traverse et à la sortie des amplificateurs opérationnels où on ne connaît pas *a priori* l'intensité du courant sortant du composant.

7. Régime sinusoïdal forcé et régime transitoire

7.1 Rappels sur les définitions

On a vu précédemment que la réponse d'un système linéaire, décrit par une équation différentielle linéaire à coefficients constants, à une excitation quelconque est la somme de la solution générale de l'équation différentielle homogène associée à l'équation différentielle (c'est-à-dire à second membre nul) et d'une solution particulière.

La solution générale de l'équation homogène associée décrit le *régime libre* tandis que la solution particulière correspond au *régime permanent*. Le régime libre a été étudié au cours de la première période. Dans ce chapitre, on s'est intéressé à la recherche d'une solution particulière dans le cas où l'excitation est sinusoïdale. On appelle *régime forcé* la réponse en régime permanent à une excitation de type sinusoïdale ou pouvant s'y ramener par décomposition en série de Fourier comme on le verra dans un chapitre ultérieur.

On rappelle que le régime libre d'un système est la réponse de ce dernier en l'absence de toute excitation.

7.2 Régime transitoire et régime forcé

L'obtention du régime forcé, comme celle de tout régime permanent, n'est pas instantanée. On observe toujours un *régime transitoire* entre deux régimes permanents qu'ils soient continus, forcés ou autres. On retiendra donc que **ces deux régimes permanent et transitoire sont toujours présents** et qu'il faut bien les distinguer.

Le régime permanent correspond à la solution particulière sinusoïdale, le régime libre à la solution générale de l'équation homogène associée. La solution générale est la superposition de ces deux régimes. Tant que le régime libre n'est pas négligeable devant le régime permanent, on parle de *régime transitoire*.

Les figures suivantes donnent le comportement du régime libre et du régime transitoire lors de l'établissement d'un régime permanent sinusoïdal pour les mêmes valeurs des composants.

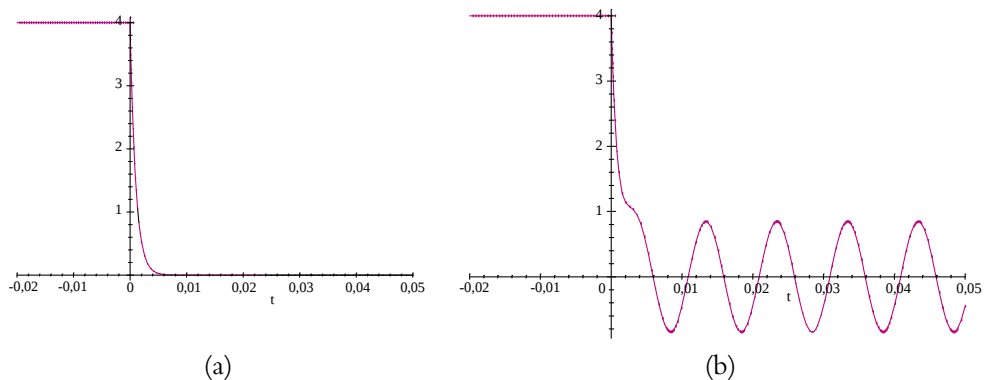


Figure 15.15 Circuit R, C avec $R = 1 \text{ k}\Omega$ et $C = 1 \text{ }\mu\text{F}$: (a) régime libre et (b) établissement d'un régime sinusoïdal d'amplitude 1 V et de fréquence 100 Hz.

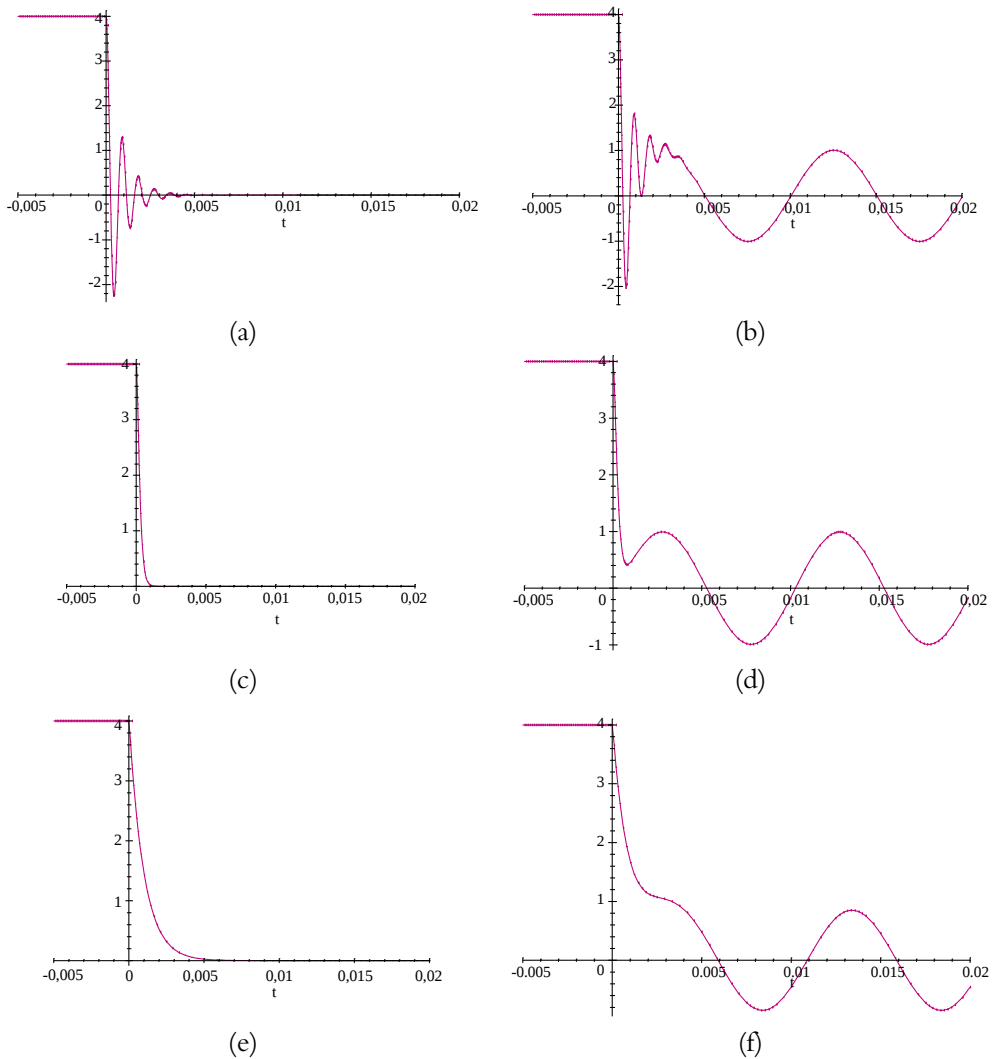


Figure 15.16 Circuit R, L, C avec $L = 20 \text{ mH}$ et $C = 1 \mu\text{F}$: (a) régime libre pour $R = 50 \Omega$, (b) établissement d'un régime sinusoïdal d'amplitude 1 V et de fréquence 100 Hz pour $R = 50 \Omega$, (c) régime libre pour $R = 283 \Omega$, (d) établissement d'un régime sinusoïdal d'amplitude 1 V et de fréquence 100 Hz pour $R = 283 \Omega$, (e) régime libre pour $R = 1 \text{ k}\Omega$, (f) établissement d'un régime sinusoïdal d'amplitude 1 V et de fréquence 100 Hz pour $R = 1 \text{ k}\Omega$.

On constate donc que le régime transitoire est très court, c'est la raison pour laquelle on peut le négliger lorsqu'on étudie le régime permanent. Il ne faut cependant jamais oublier son existence.

7.3 Importance du régime transitoire : exemple des tubes fluorescents

Les tubes d'éclairage courants sont des tubes fluorescents contenant un gaz qui n'est conducteur que lorsqu'il est ionisé. La tension nécessaire à l'ionisation du gaz (600 V à 700 V) est supérieure à la tension sinusoïdale d'alimentation (240 V). On constate que l'allumage de ces tubes n'est pas immédiat comme c'est le cas pour les ampoules à filament. Ceci s'explique par le régime transitoire précédent le régime permanent à savoir l'éclairage par les tubes.

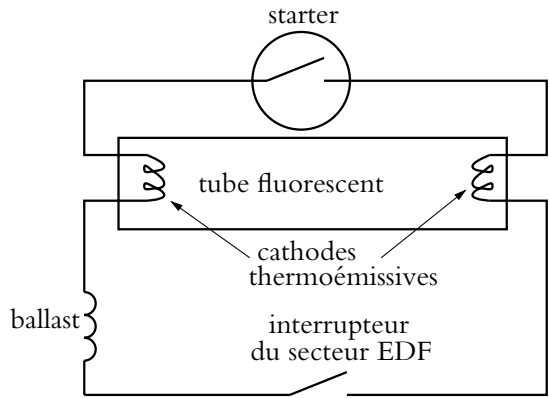


Figure 15.17 Modèle du tube fluorescent.

Le tube fluorescent peut être modélisé par la figure 15.17. Il est constitué du tube, d'une bobine appelée ballast, du tube d'éclairage et d'un interrupteur automatique appelé starter.

L'allumage du tube se déroule suivant la séquence suivante.

1. À la fermeture du circuit (lorsqu'on ferme l'interrupteur pour allumer la lumière), le starter se ferme quasiment immédiatement.
2. Après quelques secondes, le starter s'ouvre automatiquement, ce qui provoque une surtension aux bornes de la bobine. La tension aux bornes du tube devient alors suffisante pour ioniser le gaz et le rendre conducteur.
3. Le circuit est alors fermé par le tube dont le gaz ionisé émet la lumière.

Ce sont les quelques secondes nécessaires à l'ouverture du starter qui explique l'allumage retardé du tube.

8. Stabilité des circuits du premier et du second ordre (PCSI)

8.1 Stabilité d'un circuit

On dit qu'un système est *stable* si au bout d'un temps plus ou moins long, son évolution est celle d'un régime permanent. Cela signifie que le régime libre disparaît au sens où son amplitude tend vers 0.

8.2 Stabilité des circuits du premier ordre

Ces systèmes sont décrits par une équation différentielle du premier ordre dont l'écriture générale est :

$$a \frac{ds}{dt} + bs = 0$$

On notera que seul le régime libre est utile pour l'analyse de la stabilité et qu'il n'est donc pas nécessaire de tenir compte de la présence d'un éventuel second membre. La solution est de la forme :

$$s = S_0 \exp\left(-\frac{b}{a}t\right)$$

Cette solution tend vers 0 à condition que a et b soient de même signe pour que l'exponentielle soit réellement décroissante. Sinon, la solution tend vers $\pm \infty$ selon le signe de S_0 , ce qui n'a pas de réel sens physique. Dans la réalité, il existe un phénomène de saturation lié aux non-linéarités qui apparaissent pour des amplitudes importantes.

Si on revient sur les exemples des circuits R, L et R, C étudiés au chapitre 3, on note que les coefficients de l'équation différentielle sont de même signe et que le système est stable. C'est tout à fait conforme à ce qu'on a trouvé lors de l'étude de la réponse à un échelon.

Un circuit du premier ordre est stable si les deux coefficients de l'équation différentielle homogène associée sont de même signe.

8.3 Stabilité des circuits du second ordre

Comme pour les circuits du premier ordre, on peut déterminer les conditions pour qu'un système décrit par une équation différentielle du second ordre soit stable. La définition est toujours la même : on cherche à obtenir un régime libre qui disparaît au bout d'un temps plus ou moins long au sens où son amplitude tend vers 0.

On s'intéresse ici aux systèmes du second ordre dont l'équation différentielle s'écrit :

$$a \frac{d^2s}{dt^2} + b \frac{ds}{dt} + cs = 0$$

L'équation caractéristique associée est :

$$ar^2 + br + c = 0$$

et son discriminant vaut : $\Delta = b^2 - 4ac$.

- Si $\Delta < 0$ alors :

$$r = \frac{-b \pm j\sqrt{4ac - b^2}}{2a}$$

On a dans ce cas des solutions sinusoïdales modulées en amplitude dont la pulsation est donnée par la partie imaginaire de la solution de l'équation caractéristique. Leur amplitude est modulée de façon exponentielle : elle est proportionnelle à $\exp\left(-\frac{b}{2a}t\right)$. Le système sera stable si l'amplitude tend vers 0 ou encore si l'argument de l'exponentielle est négatif. Ce sera le cas si a et b sont de même signe.

- Si $\Delta = 0$ alors :

$$r = -\frac{b}{2a}$$

et on obtient la même condition que dans le cas précédent pour les mêmes raisons.

- Si $\Delta \geq 0$ alors :

$$r = \frac{-b \pm \sqrt{b^2 - 4ac}}{2a}$$

Les solutions de l'équation différentielle sont donc des exponentielles qui doivent tendre vers 0 pour assurer la stabilité du système : l'argument des exponentielles doit être négatif.

- Si a et c sont de même signe, $\sqrt{b^2 - 4ac} \leq b$ et le numérateur est du signe opposé à b .
 - ◇ Si a et b sont de même signe, $r < 0$ et le système est stable,
 - ◇ Si a et b sont de signe opposé, $r > 0$ et le système est instable.
- Si a et c sont de signe opposé, alors $-b + \sqrt{b^2 - 4ac} > 0$ et le système est instable.

Un système du second ordre est stable lorsque tous les coefficients de l'équation différentielle homogène associée sont de même signe.

On remarquera que si les coefficients ne sont pas tous de même signe, il est possible que le système soit stable à condition que la recherche des constantes d'intégration conduise à obtenir un coefficient nul pour la fonction instable.

A. Applications directes du cours

1. Équivalence de deux dipôles

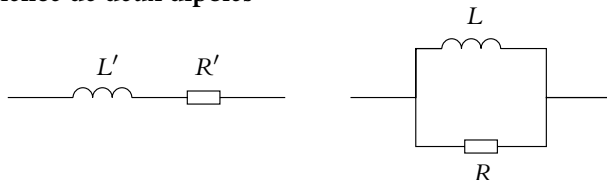


Figure 15.18

On travaille en régime sinusoïdal de pulsation ω .

1. Pour quelles valeurs de R' et L' , a-t-on l'équivalence des montages série et parallèle représentés ci-dessus ?
2. Pour quelle(s) valeur(s) de la fréquence a-t-on $\frac{L}{R} = \frac{L'}{R'}$?

2. Lois générales en régime sinusoïdal

Soit le montage :

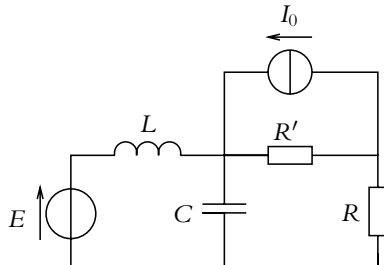


Figure 15.19

Déterminer le courant parcourant la résistance R en appliquant :

1. les équivalences entre modèles de Thévenin et de Norton,
2. le théorème de Millman.

3. Courant et tension en phase

On veut que le courant parcourant R_1 soit en phase avec la tension aux bornes du générateur.

Quelle est la condition à vérifier pour qu'il en soit ainsi en régime sinusoïdal ?

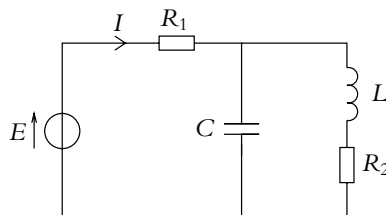


Figure 15.20

B. Exercices et problèmes

1. Étude de l'impédance d'un quartz piezoélectrique, d'après ENSTIM 2004

Le schéma électrique simplifié d'un quartz est donné sur la figure (15.21). On néglige sa résistance R . On donne $L = 500 \text{ mH}$; $C_S = 0,08 \text{ pF}$ et $C_P = 8 \text{ pF}$.

1. a) Calculer l'impédance complexe du quartz vue entre les bornes A et B et la mettre sous la forme :

$$\underline{Z}_{AB} = \left(-\frac{j}{\alpha\omega} \right) \frac{1 - \frac{\omega^2}{\omega_r^2}}{1 - \frac{\omega^2}{\omega_a^2}}$$

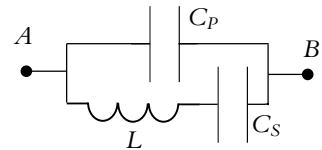


Figure 15.21

où j est le nombre complexe tel que $j^2 = -1$, et α , ω_r et ω_a sont à déterminer.

Montrer aussi que $\omega_a^2 > \omega_r^2$

b) Calculer les valeurs numériques des fréquences f_a et f_r correspondant aux pulsations ω_a et ω_r .

2. a) Étudier le comportement inductif (partie imaginaire positive) ou capacitif (partie imaginaire négative) du quartz en fonction de la fréquence. Représenter l'argument de $\underline{Z}_{AB}$ en fonction de ω .

b) Tracer l'allure de $Z_{AB} = |\underline{Z}_{AB}|$, module de l'impédance complexe du quartz, en fonction de la fréquence.

c) Comment est modifiée cette courbe si on tient compte de la résistance de la bobine ?

2. Adaptation d'impédance, d'après Polytechnique 2004

Un dipôle électrocinétique linéaire passif est, en régime sinusoïdal permanent, caractérisé par son impédance complexe $\underline{Z} = R + jX$.

Le système étudié (réacteur à plasma) est modélisé par un circuit série $R_P - C_P$. On veut diminuer au maximum la partie imaginaire (appelée *partie réactive*) de cette impédance $\underline{Z}_P$. Pour cela, on réalise le circuit de la figure (15.22).

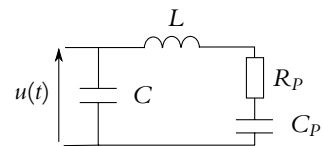


Figure 15.22

1. Exprimer l'admittance totale $\underline{Y}$ du dipôle de la figure (15.22). Déterminer l'expression de C qui annule la partie réactive de $\underline{Z}_P$.

2. La condition précédente étant réalisée, déterminer l'expression de l'impédance $\underline{Z}$ totale du dipôle, notée alors R_1 .

3. Impédance itérative, d'après ICNA 2000

Le circuit de la figure (15.23) est alimenté entre ses bornes "d'entrée" A et B par un générateur qui délivre à l'instant t la tension $u_e(t)$. Cette tension, sinusoïdale, a pour amplitude complexe $\underline{U}_e$, et pour pulsation ω .

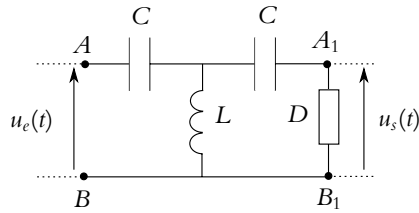


Figure 15.23

"En sortie", entre les bornes A_1 et B_1 , est placé un dipôle D d'impédance complexe $\underline{Z}$.

1. Calculer en fonction de L , C , ω et $\underline{Z}$ l'impédance complexe $\underline{Z}_e$ du circuit vue entre les bornes A et B (impédance d'entrée).
2. En déduire la valeur de $\underline{Z}$ pour laquelle $\underline{Z}_e = \underline{Z}$ (appelée alors impédance itérative).
3. a) On donne $L = 1$ mH et $C = 0,2$ μ F. Montrer que pour les hautes fréquences (estimer une limite inférieure), on peut considérer que la ligne est purement résistive (équivalente à une résistance R)
b) Calculer la valeur de la résistance R .
4. Examiner le comportement du circuit pour $\omega = 0$ et $\omega \rightarrow +\infty$.

4. Couplage par induction

Les deux circuits sont couplés par induction avec M comme coefficient d'induction mutuelle, ce qui signifie que la bobine du circuit α induit une tension $M \frac{di_\alpha}{dt}$ aux bornes de la bobine du circuit β .

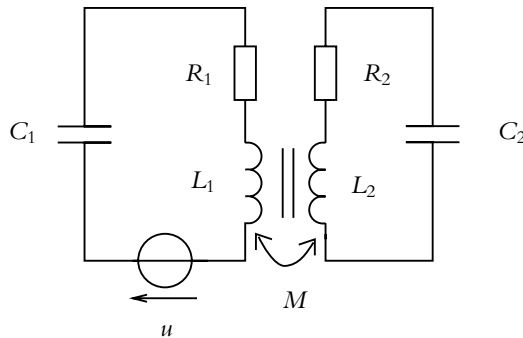


Figure 15.24

1. Écrire les équations différentielles relatives aux charges q_1 et q_2 dans chaque circuit.
2. On cherche une solution sous la forme de charge oscillant à la même pulsation ω dans chaque circuit. En déduire le système de deux équations à deux inconnues permettant d'obtenir l'expression complexe des charges complexes associées.

3. On pose $X = L\omega - \frac{1}{C\omega}$. Déterminer l'expression des courants sous forme complexe en fonction de M , ω , X , R et l'amplitude U de la source de tension.
4. En déduire l'amplitude des courants.

5. Étude d'un circuit en régime sinusoïdal, d'après ICARE PSI/PT 1999

1. On considère le dipôle constitué d'une résistance R en parallèle avec une capacité C . Déterminer la résistance R' et la capacité C' qui, en série, ont la même impédance que ce dipôle pour une pulsation ω donnée de la tension appliquée.

2. Tracer sur un même graphique les courbes représentatives de $\frac{R}{R'}$ et $\frac{C}{C'}$ en fonction du rapport $\frac{\omega}{\omega_0}$ où $\omega_0 = \frac{1}{RC}$.

3. On considère désormais le dispositif constitué de l'association en série des dipôles formés d'une résistance R et d'une capacité C respectivement associées en série et en parallèle. On alimente l'ensemble par une tension sinusoïdale $e(t)$ d'amplitude E , de pulsation ω et de phase initiale nulle.

Déterminer le rapport $\frac{u}{e}$ où $u(t)$ désigne la tension aux bornes de l'association en parallèle de R et C . Même question pour $v(t)$ la tension aux bornes de l'association en série de R et C .

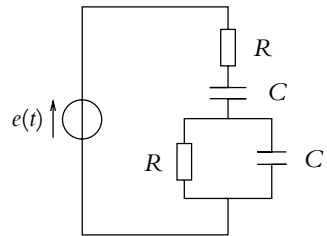


Figure 15.25

4. Exprimer l'amplitude de u .
5. Tracer et justifier l'allure du rapport entre l'amplitude de u et celle de e en fonction de la pulsation ω . On donnera toutes les valeurs particulières.
6. Exprimer le déphasage de u par rapport à e .
7. Tracer et justifier l'allure du déphasage de u par rapport à e en fonction de la pulsation ω .
8. Déterminer l'intensité complexe circulant dans le générateur.
9. Étudier le déphasage de l'intensité par rapport à e en fonction de la pulsation ω .
10. On branche les deux bornes extrêmes d'un potentiomètre en parallèle avec le générateur et on relie la borne intermédiaire via un voltmètre à la connexion entre les associations en série et en parallèle de R et C . On souhaite que le voltmètre indique une tension nulle. Pourquoi choisit-on de se placer à la pulsation ω_0 ?
11. Déterminer la position qu'il faut donner au potentiomètre pour que le voltmètre indique une tension nulle.

16

Circuit R, L, C série en régime sinusoïdal forcé et résonances

1. Circuit R, L, C série soumis à une excitation sinusoïdale

1.1 Différentes tensions et montages correspondants

Le circuit étudié dans ce chapitre est constitué d'un résistor de résistance R , d'une bobine d'inductance L et d'un condensateur de capacité C associées en série. Suivant le dipôle aux bornes duquel on étudiera la tension, on rencontrera plusieurs types de comportements.

a) Tension aux bornes du condensateur

Aux bornes du condensateur, on aura une image de sa charge ou de la primitive de l'intensité :

$$u_C = \frac{q}{C} \quad \text{ou} \quad u_C = \frac{1}{C} \int i dt$$

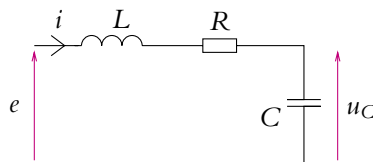


Figure 16.1 Tension aux bornes de C .

b) Intensité ou tension aux bornes de la résistance

Aux bornes du résistor, on aura une image de l'intensité :

$$u_R = Ri$$

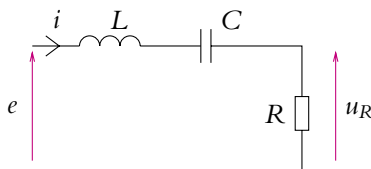


Figure 16.2 Tension aux bornes de R .

c) Tension aux bornes de la bobine

Aux bornes de la bobine dont on néglige la résistance interne, on aura une image de la dérivée de l'intensité :

$$u_L = L \frac{di}{dt}$$

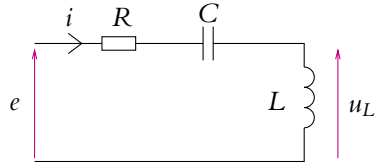


Figure 16.3 Tension aux bornes de L .

On aura l'une ou l'autre des situations suivant les positions relatives des composants et notamment l'endroit où la masse du circuit sera imposée par le G.B.F. ou l'oscilloscope généralement reliés à la terre (Cf. chapitre sur l'instrumentation électrique). Ainsi le dipôle aux bornes duquel on veut étudier la tension devra avoir une borne commune avec le générateur, cette borne sera la masse du circuit (le plus souvent reliée à la terre).

1.2 Équation différentielle

a) Relative à la tension aux bornes du condensateur

L'écriture de la loi des mailles dans le circuit donne :

$$e = L \frac{di}{dt} + Ri + u_C \quad (16.1)$$

Or la relation entre intensité et tension aux bornes d'une capacité est : $i = C \frac{du_C}{dt}$, ce qui permet d'en déduire l'équation différentielle vérifiée par u_C aux bornes du condensateur :

$$e = LC \frac{d^2 u_C}{dt^2} + RC \frac{du_C}{dt} + u_C$$

ou encore :

$$\frac{d^2 u_C}{dt^2} + \frac{R}{L} \frac{du_C}{dt} + \frac{u_C}{LC} = \frac{e}{LC} \quad (16.2)$$

b) Relative à l'intensité ou à la tension aux bornes de la résistance

L'intensité s'obtient par dérivation du u_C et utilisation de la relation $i = C \frac{du_C}{dt}$. On dérive donc l'équation différentielle obtenue précédemment par rapport au temps :

$$\frac{d^3 u_C}{dt^3} + \frac{R}{L} \frac{d^2 u_C}{dt^2} + \frac{1}{LC} \frac{du_C}{dt} = \frac{1}{LC} \frac{de}{dt}$$

soit en remplaçant $\frac{du_C}{dt}$ par $\frac{i}{C}$:

$$\frac{d^2 i}{dt^2} + \frac{R}{L} \frac{di}{dt} + \frac{i}{LC} = \frac{1}{L} \frac{de}{dt} \quad (16.3)$$

c) Relative à la tension aux bornes de la bobine

La tension aux bornes de la bobine s'obtient en dérivant l'intensité i et en la multipliant par L : $u_L = L \frac{di}{dt}$. En procédant comme au paragraphe précédent, on en déduit l'équation différentielle vérifiée par u_L :

$$\frac{d^2 u_L}{dt^2} + \frac{R}{L} \frac{du_L}{dt} + \frac{u_L}{LC} = \frac{d^2 e}{dt^2}$$

1.3 Régime transitoire et régime forcé

On a déjà montré au chapitre précédent (Cf. figures 15.16) que la réponse du circuit à une excitation sinusoïdale est la somme de la solution générale de l'équation différentielle homogène associée à l'équation différentielle et d'une solution particulière. La première décrit le régime libre étudié en première période et la seconde le régime permanent ou régime sinusoïdal forcé qui subsiste quand le régime libre s'est annulé.

1.4 Passage à la notation complexe

Pour la recherche du régime sinusoïdal forcé, on a vu au chapitre précédent l'intérêt de l'utilisation de la notation complexe : cette approche permet de remplacer la résolution d'une équation différentielle par celle d'une équation algébrique sur le corps des complexes. Il s'agit d'une simplification importante pour la résolution du problème.

a) Expression de l'intensité et de la tension aux bornes de la résistance

En notation complexe, l'équation différentielle (16.3) s'écrit :

$$(j\omega)^2 \underline{i} + \frac{R}{L} j\omega \underline{i} + \frac{\underline{i}}{LC} = \frac{1}{L} j\omega \underline{e} \quad \text{soit} \quad \underline{i} = \frac{\underline{e}}{R + j\left(L\omega - \frac{1}{C\omega}\right)}$$



On obtient directement cette équation en traduisant l'équation (16.1) en notations complexes grâce aux impédances complexes.

On a l'habitude d'introduire la pulsation ω_0 et le nombre sans dimension Q définis par $\omega_0 = \frac{1}{\sqrt{LC}}$ et $Q = \frac{1}{R} \sqrt{\frac{L}{C}}$. L'expression de l'intensité complexe devient alors :

$$\underline{i} = \frac{\underline{e}}{R \left(1 + jQ \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right) \right)}$$

On en déduit l'expression de la tension complexe aux bornes de R :

$$\underline{u_R} = R\underline{i} = \frac{R\underline{e}}{R + j\left(L\omega - \frac{1}{C\omega}\right)} = \frac{\underline{e}}{1 + jQ\left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega}\right)}$$

b) Expression de la tension aux bornes du condensateur

La tension aux bornes du condensateur s'obtient alors facilement par :

$$\underline{u_C} = \frac{\underline{i}}{jC\omega} = \frac{\underline{e}}{(1 - LC\omega^2) + jRC\omega}$$

soit en fonction de Q et ω_0 :

$$\underline{u_C} = \frac{\underline{e}}{1 - \frac{\omega^2}{\omega_0^2} + j\frac{1}{Q}\frac{\omega}{\omega_0}}$$

c) Expression de la tension aux bornes de la bobine

De même, la tension aux bornes de la bobine est :

$$\begin{aligned}\underline{u_L} &= jL\omega\underline{i} = \frac{jL\omega\underline{e}}{R + j\left(L\omega - \frac{1}{C\omega}\right)} \\ &= \frac{-LC\omega^2\underline{e}}{1 - LC\omega^2 + jRC\omega} = \frac{-\frac{\omega^2}{\omega_0^2}\underline{e}}{1 - \frac{\omega^2}{\omega_0^2} + j\frac{1}{Q}\frac{\omega}{\omega_0}}\end{aligned}$$

2. Résonance en intensité

On étudie l'intensité en fonction de la fréquence. Son expression complexe est :

$$\underline{i} = \frac{\underline{e}}{R + j\left(L\omega - \frac{1}{C\omega}\right)} = \frac{\underline{e}}{R\left(1 + jQ\left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega}\right)\right)}$$

2.1 Étude de l'amplitude

L'amplitude est le module de l'intensité complexe soit :

$$I = |\underline{i}| = \frac{E}{\sqrt{R^2 + \left(L\omega - \frac{1}{C\omega}\right)^2}} = \frac{E}{R\sqrt{1 + Q^2\left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega}\right)^2}}$$

en notant E l'amplitude de la tension sinusoïdale imposée $e(t)$.

On étudie les variations de I en fonction de la fréquence f ou de la pulsation $\omega = 2\pi f$.

On remarque que $x \rightarrow \frac{1}{\sqrt{x}}$ est une fonction strictement décroissante donc les varia-

tions de I sont opposées à celles de $f(\omega) = R^2 + \left(L\omega - \frac{1}{C\omega}\right)^2$.

On peut remarquer que $\forall \omega, f(\omega) \geq R^2$ et que :

- quand ω tend vers 0, $f(\omega)$ tend vers $+\infty$ comme $\frac{1}{C^2\omega^2}$ qui est le terme prépondérant en 0 ;
- quand ω tend vers $+\infty$, $f(\omega)$ tend vers $+\infty$ comme $L^2\omega^2$ qui est le terme prépondérant à l'infini.
- $f\left(\omega_0 = \frac{1}{\sqrt{LC}}\right) = R^2$.

La fonction f passe donc par un minimum en $\omega_0 = \frac{1}{\sqrt{LC}}$.

On en déduit donc que :

- l'amplitude I de l'intensité admet un maximum qui vaut $\frac{E}{R}$ pour la pulsation $\omega = \omega_0$; on dit qu'il y a *résonance en intensité*. En effet, par définition, le phénomène de résonance correspond à l'existence d'un maximum de la quantité étudiée en fonction de la fréquence (ou de la pulsation).

ω_0 est appelée *pulsation de résonance* en intensité ou pulsation propre du circuit.

- L'amplitude I de l'intensité croît pour $\omega < \omega_0$ et décroît ensuite.

Les valeurs limites sont les suivantes :

- quand ω tend vers 0, $f(\omega) \sim \frac{1}{C^2\omega^2}$ donc $I \sim C\omega E$ et tend vers 0 ; la tangente à l'origine a une pente non nulle, égale à CE ;
- quand ω tend vers $+\infty$, l'amplitude I de l'intensité tend vers 0 ;
- quand $\omega = \omega_0$, $I = \frac{E}{R}$.

2.2 Étude du déphasage

Le déphasage φ entre $i(t)$ et $e(t)$ est :

$$\varphi = \text{Arg}(\hat{i}) - \text{Arg}(\hat{e}) = \text{Arg}(1/R) - \text{Arg}\left(1 + jQ\left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega}\right)\right)$$

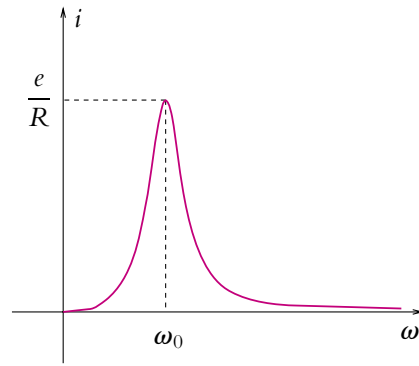


Figure 16.4 Résonance en intensité.

$$\varphi = -\text{Arg} \left(1 + jQ \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right) \right) = -\psi$$

où $\psi = \text{Arg } \underline{D}$ avec $\underline{D} = 1 + jQ \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)$.

La détermination de l'argument s'obtient en calculant sa tangente et en précisant l'intervalle de longueur π auquel il appartient. Le deuxième point est fondamental : la tangente étant π -périodique, on a une incertitude de π que la connaissance de l'intervalle auquel appartient l'argument permet de lever. La détermination de cet intervalle se fait en étudiant les signes du cosinus et/ou du sinus de l'argument.

Dans le cas présent :

$$\tan \varphi = -\tan \psi = -Q \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)$$

La partie réelle de $\underline{D}$ est positive (alors que sa partie imaginaire change de signe suivant les valeurs de ω) donc $\cos \psi$ est positif et $\varphi \in \left[-\frac{\pi}{2}, \frac{\pi}{2}\right]$.

On peut établir que le déphasage est une fonction décroissante de ω .

D'autre part, les valeurs limites sont les suivantes :

- quand ω tend vers 0, $\tan \varphi$ tend vers $+\infty$ donc comme $\varphi \in \left[-\frac{\pi}{2}, \frac{\pi}{2}\right]$, φ tend vers $\frac{\pi}{2}$;
- quand ω tend vers $+\infty$, $\tan \varphi$ tend vers $-\infty$ donc comme $\varphi \in \left[-\frac{\pi}{2}, \frac{\pi}{2}\right]$, φ tend vers $-\frac{\pi}{2}$;
- quand $\omega = \omega_0$, $\tan \varphi$ tend vers 0 donc comme $\varphi \in \left[-\frac{\pi}{2}, \frac{\pi}{2}\right]$, φ tend vers 0.

Cette étude permet de donner l'allure du déphasage φ entre l'intensité et la tension d'entrée en fonction de la pulsation ω :

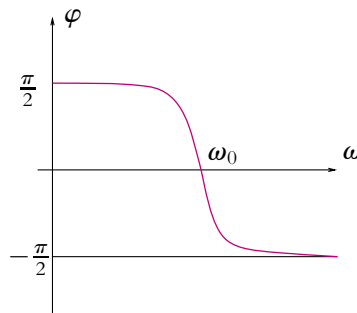


Figure 16.5 Déphasage de la résonance en intensité.

Dans la suite, on n'utilisera plus que les expressions en fonction de Q et de ω_0 , la pulsation propre.

2.3 Bande passante à -3 dB

On définit l'intensité en décibels¹ par :

$$i_{\text{dB}} = 20 \log_{10} \frac{|i|}{|i_{\text{ref}}|}$$

où $|i_{\text{ref}}|$ est une intensité de référence, que l'on peut prendre égale à 1 ampère.

On peut alors définir la bande passante à -3 dB comme l'ensemble des fréquences ou des pulsations telles que

$$|i(\omega)| \geq \frac{i_{\text{max}}}{\sqrt{2}} \quad (16.4)$$

Donc

$$|i(\omega)| \geq \frac{i_{\text{max}}}{\sqrt{2}} \Leftrightarrow i_{\text{dB}} \geq 20 \log_{10} \frac{i_{\text{max}}}{|i_{\text{ref}}|} - 20 \log_{10} \sqrt{2} = i_{\text{dB,max}} - 3$$

car $20 \log_{10} \sqrt{2} \simeq 3,0$ dB. Ceci explique le nom de bande passante à -3 dB.

Pour la détermination de la bande passante, on utilisera plutôt la définition (16.4) tout en parlant de bande passante à -3 dB. Ici $i_{\text{max}} = \frac{E}{R}$, obtenue lorsque $\omega = \omega_0$.

On cherche donc les pulsations ω telles que :

$$\frac{E}{R \sqrt{1 + Q^2 \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)^2}} \geq \frac{E}{\sqrt{2} R} \quad (16.5)$$

soit :

$$\begin{aligned} 1 + Q^2 \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)^2 &\leq 2 \\ \Leftrightarrow Q^2 \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)^2 - 1 &\leq 0 \\ \Leftrightarrow (Q(\omega^2 - \omega_0^2) - \omega_0 \omega) (Q(\omega^2 - \omega_0^2) + \omega_0 \omega) &\leq 0 \end{aligned}$$



On remarque qu'il ne faut surtout pas développer le carré mais au contraire factoriser l'expression.

D'après l'allure de la courbe $I(\omega)$, il suffit de chercher les pulsations ω_1 et ω_2 ($> \omega_1$) pour lesquelles l'égalité est satisfaite. La bande passante sera $[\omega_1, \omega_2]$.

¹ L'étude des grandeurs en décibels sera faite dans le chapitre sur les filtres.

On résout tout d'abord :

$$Q\omega^2 - \omega_0\omega - Q\omega_0^2 = 0$$

Le discriminant vaut : $\Delta = \omega_0^2 + 4Q^2\omega_0^2 > 0$, on a donc deux solutions réelles :

$$\omega_{1,1} = \frac{\omega_0 - \omega_0\sqrt{1+4Q^2}}{2Q} \quad \text{et} \quad \omega_{1,2} = \frac{\omega_0 + \omega_0\sqrt{1+4Q^2}}{2Q}$$

Seules les pulsations positives sont physiquement acceptables, on ne gardera donc que la solution $\omega_{1,2} = \omega_1$.

De même, on résout :

$$Q\omega^2 + \omega_0\omega - Q\omega_0^2 = 0$$

Les solutions de cette équation sont :

$$\omega_{2,1} = \frac{-\omega_0 - \omega_0\sqrt{1+4Q^2}}{2Q} \quad \text{et} \quad \omega_{2,2} = \frac{-\omega_0 + \omega_0\sqrt{1+4Q^2}}{2Q}$$

On ne garde que la solution positive $\omega_{2,2} = \omega_2$.

Les limites de la bande passante sont donc :

$$\omega_1 = \frac{-\omega_0 + \omega_0\sqrt{1+4Q^2}}{2Q} \quad \text{et} \quad \omega_2 = \frac{\omega_0 + \omega_0\sqrt{1+4Q^2}}{2Q}$$

La largeur de la bande passante est alors donnée par :

$$\Delta\omega = \omega_2 - \omega_1 = \frac{\omega_0}{Q} = \frac{L}{R}$$

On dit que la résonance est aiguë si la bande passante est faible et qu'elle est floue dans le cas inverse.

2.4 Déphasage et bande passante

À la résonance, on a $\omega = \omega_0$ soit un déphasage nul d'après les études précédentes. Cela fournit donc un moyen facile pour déterminer expérimentalement la pulsation de résonance : au lieu de chercher la valeur donnant un maximum de l'amplitude de l'intensité, on cherche la valeur donnant un déphasage nul, par exemple en observant une droite en représentation de Lissajous (Cf. chapitre sur l'instrumentation pour la description de la méthode).



La définition de la résonance n'est pas un déphasage nul mais un maximum de la quantité étudiée. Ce n'est donc pas parce que, dans le cas de la résonance en intensité, on a comme résultat remarquable que le déphasage est nul qu'il faudra en déduire qu'il s'agit de la définition de la résonance.

Pour chacune des limites de la bande passante, on a :

- pour ω_1 : $\tan \varphi_1 = Q \left(\frac{\omega_1}{\omega_0} - \frac{\omega_0}{\omega_1} \right) = 1$ soit $\varphi_1 = \frac{\pi}{4}$,
- pour ω_2 : $\tan \varphi_2 = Q \left(\frac{\omega_2}{\omega_0} - \frac{\omega_0}{\omega_2} \right) = -1$ soit $\varphi_2 = -\frac{\pi}{4}$.

On dit que la bande passante peut également être définie par des fréquences (ou des pulsations) quadrantales, c'est-à-dire définissant une quadrature de phase entre les deux grandeurs. Cela peut également fournir une méthode pour déterminer expérimentalement les pulsations ω_1 et ω_2 . Elle est généralement moins pratique que celle permettant la mesure de ω_0 sauf dans le cas où l'oscilloscope numérique affiche les valeurs du déphasage entre les deux signaux.

2.5 Facteur de qualité et bande passante

La relation $Q = \frac{\omega_0}{\Delta\omega}$ peut constituer une définition du facteur de qualité. En remplaçant ω_0 par $\frac{1}{\sqrt{LC}}$ et $\Delta\omega$ par $\frac{L}{R}$, on retrouve les différentes expressions de Q :

$$Q = \frac{L\omega_0}{R} = \frac{1}{RC\omega_0} = \frac{1}{R} \sqrt{\frac{L}{C}}$$

La résonance sera donc d'autant plus aiguë que la bande passante sera faible et par conséquent le facteur de qualité grand.

Elle sera d'autant plus floue que la bande passante sera grande et par conséquent le facteur de qualité faible.

2.6 Facteur de qualité ou coefficient de surtension

Il existe une interprétation physique du facteur de qualité.

On se place dans les conditions de la résonance en intensité, à savoir $\omega = \omega_0$, et on s'intéresse à la tension aux bornes de la capacité ou de l'inductance.

Aux bornes de la capacité, la tension vaut :

$$\underline{u_C} = \frac{\underline{i}}{jC\omega}$$

Or, pour $\omega = \omega_0$, $\underline{i} = \frac{E}{R}$ d'où

$$|\underline{u_C}|(\omega_0) = \frac{E}{RC\omega_0}$$

De même, aux bornes de l'inductance :

$$\underline{u_L} = jL\omega \underline{i}$$

donc

$$|u_L|(\omega_0) = \frac{L\omega_0 E}{R}$$

Comme $L\omega_0 = \frac{1}{C\omega_0}$, ces deux tensions sont égales :

$$|u_L|(\omega_0) = |u_C|(\omega_0) = \frac{E}{RC\omega_0} = \frac{L\omega_0 E}{R}$$

Le facteur de qualité correspond alors au facteur de surtension défini par :

$$Q = \frac{|u_L| \text{ à la résonance en intensité }}{E} = \frac{|u_C| \text{ à la résonance en intensité }}{E} = \frac{1}{RC\omega_0} = \frac{L\omega_0}{R}$$

Ce coefficient peut prendre des valeurs importantes pour un « bon » circuit. La tension aux bornes de la capacité et celle aux bornes de l'inductance peuvent devenir importantes. On parle d'effet de surtension, d'où le nom donné à ce coefficient. On peut cependant noter qu'il est nécessaire, pour ne pas endommager le composant, que la tension aux bornes de la capacité reste inférieure à la tension de claquage du composant. Il est donc fondamental de procéder à un choix judicieux des valeurs des composants en tenant compte de la tension maximale délivrée par le G.B.F.

Le facteur de qualité est donc une mesure :

- de la surtension aux bornes de la capacité et de l'inductance à la résonance,
- de l'acuité de la résonance,

et ce uniquement pour la résonance **en intensité**, comme on le verra dans le paragraphe suivant.

Ce qu'il faut retenir :

Le phénomène de résonance en intensité correspond à l'obtention d'un maximum d'intensité en fonction de la fréquence ou à un déphasage nul entre l'intensité et la tension fournie par le générateur.

Cette résonance existe toujours quels que soient les paramètres R , C et L du circuit.

La bande passante correspond aux fréquences pour lesquelles l'intensité est supérieure à $\frac{i_{\max}}{\sqrt{2}}$ ou bien pour lesquelles l'intensité en décibels est supérieure à $i_{\text{dB,max}} - 3$, ce qui justifie la dénomination de bande passante à -3 dB. On reviendra sur ce point dans un chapitre ultérieur.

Le facteur de qualité peut être défini comme coefficient de surtension aux bornes de la capacité ou aux bornes de l'inductance à résonance ou bien comme une mesure de l'acuité de la bande passante. Il est important de se rappeler que le facteur de qualité est un nombre sans dimension : ce sera le rapport de deux quantités homogènes que sont ou bien la pulsation propre et la bande passante de la résonance en intensité, ou

bien deux tensions. D'autre part, un bon facteur de qualité a une valeur importante et correspond à une bande passante étroite : la bande passante sera donc au dénominateur dans l'expression de Q .

3. Résonance aux bornes de la capacité

La tension aux bornes de la capacité vaut :

$$\underline{u}_C = \frac{\underline{i}}{jC\omega} = \frac{\underline{\varepsilon}}{jC\omega \left(R + j\left(L\omega - \frac{1}{C\omega}\right)\right)} = \frac{\underline{\varepsilon}}{(1 - LC\omega^2) + jRC\omega}$$

soit en fonction de ω_0 et Q :

$$\underline{u}_C = \frac{\underline{\varepsilon}}{1 - \frac{\omega^2}{\omega_0^2} + j\frac{1}{Q}\frac{\omega}{\omega_0}}$$

3.1 Étude de l'amplitude

$$|\underline{u}_C| = \frac{E}{\sqrt{\left(1 - \frac{\omega^2}{\omega_0^2}\right)^2 + \frac{1}{Q^2}\frac{\omega^2}{\omega_0^2}}}$$

Pour étudier les variations de $|\underline{u}_C|$ en fonction de la fréquence, on étudie celles de :

$$f(\omega) = \left(1 - \frac{\omega^2}{\omega_0^2}\right)^2 + \frac{1}{Q^2}\frac{\omega^2}{\omega_0^2}$$

La fonction $x \rightarrow \frac{1}{\sqrt{x}}$ étant strictement décroissante, les variations de $|\underline{u}_C|$ seront inverses de celles de f .

$$f'(\omega) = 2\left(-\frac{2\omega}{\omega_0^2}\right)\left(1 - \frac{\omega^2}{\omega_0^2}\right) + 2\frac{\omega}{Q^2\omega_0^2} = \frac{2\omega}{\omega_0^2}\left(\frac{1}{Q^2} - 2 + 2\frac{\omega^2}{\omega_0^2}\right)$$

La fonction $f'(\omega)$ s'annule pour $\omega = 0$ et pour $2\frac{\omega^2}{\omega_0^2} = 2 - \frac{1}{Q^2}$. Cette deuxième solution n'existe que si $2 > \frac{1}{Q^2}$ soit $Q > \frac{1}{\sqrt{2}}$.

On aura alors à distinguer deux cas :

1. si $Q < \frac{1}{\sqrt{2}}$ alors $f'(\omega) \geq 0$ pour toute valeur de $\omega > 0$, f est croissante et $|\underline{u}_C|$ décroissante ;
2. si $Q > \frac{1}{\sqrt{2}}$ alors $f'(\omega) \geq 0$ pour $\omega \geq \omega_r = \omega_0\sqrt{1 - \frac{1}{2Q^2}}$, la fonction f est décroissante de $\omega = 0$ à ω_r . On note que $\omega_r < \omega_0$.

On aura donc un maximum et un phénomène de résonance uniquement si $Q > \frac{1}{\sqrt{2}}$.

L'étude des valeurs limites donne :

- quand $\omega \rightarrow 0$, $|\underline{u}_C| \rightarrow E$,
- quand $\omega \rightarrow +\infty$, $|\underline{u}_C| \rightarrow 0$,
- quand $\omega \rightarrow \omega_r = \omega_0 \sqrt{1 - \frac{1}{2Q^2}}$ (lorsqu'elle existe) :

$$|\underline{u}_C| = \frac{E}{\sqrt{\frac{1}{Q^2} - \frac{1}{4Q^4}}} = \frac{2Q^2 E}{\sqrt{4Q^2 - 1}}$$

en fonction du facteur de qualité².

On remarque que la résonance aux bornes du condensateur n'existe pas toujours contrairement à la résonance en intensité. D'autre part, la fréquence de résonance aux bornes du condensateur n'est pas la même que celle obtenue pour la résonance en intensité.

3.2 Étude du déphasage

On appelle φ_C l'argument de $\underline{u}_C$, c'est donc le déphasage de $u_c(t)$ par rapport à $e(t)$.

$$\begin{cases} \tan \varphi_C = \tan \left(-\text{Arg} \left(1 - \frac{\omega^2}{\omega_0^2} + j \frac{1}{Q} \frac{\omega}{\omega_0} \right) \right) \\ \sin \varphi_C < 0 \end{cases}$$

soit

$$\begin{cases} \tan \varphi_C = -\frac{\frac{1}{Q} \frac{\omega}{\omega_0}}{1 - \frac{\omega^2}{\omega_0^2}} \\ \varphi_C \in [-\pi, 0] \end{cases}$$

On peut établir que l'argument de $\underline{u}_C$ décroît.

Les valeurs limites sont les suivantes :

- quand $\omega \rightarrow 0$, $\tan \varphi_C \rightarrow 0^-$ avec $\varphi_C \in [-\pi, 0]$ donc $\varphi_C \rightarrow 0$,
- quand $\omega \rightarrow +\infty$, $\tan \varphi_C \rightarrow 0^+$ avec $\varphi_C \in [-\pi, 0]$ donc $\varphi_C \rightarrow -\pi$,
- quand $\omega \rightarrow \omega_0$, $\tan \varphi_C \rightarrow -\infty$ avec $\varphi_C \in [-\pi, 0]$ donc $\varphi_C \rightarrow -\frac{\pi}{2}$.

² On introduit parfois le paramètre ξ tel que $\xi = \frac{1}{2Q}$, le maximum s'exprime alors

$$|\underline{u}_C| = \frac{E}{2\xi \sqrt{1 - \xi^2}}$$

Des études précédentes, on déduit :

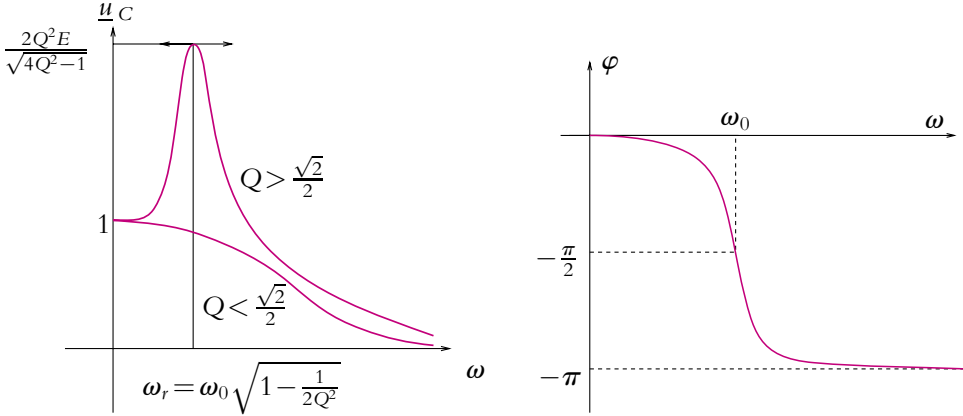


Figure 16.6 Amplitude et déphasage de la résonance aux bornes de C .

On remarque que $\text{Arg}(\underline{u}_C) = \text{Arg}(\underline{u}_R) - \frac{\pi}{2}$.

3.3 Bande passante à -3 dB

La bande passante est toujours définie comme la zone de fréquences pour laquelle la tension est supérieure à sa valeur maximale divisée par $\sqrt{2}$:

$$|\underline{u}_C(\omega)| \geq \frac{|\underline{u}_{C,\max}|}{\sqrt{2}}$$

- Si $Q < \frac{\sqrt{2}}{2}$:

Le module de $\underline{u}_C$ est décroissant donc la bande passante sera définie par $\omega < \omega_1$.

On cherche ω_1 telle que :

$$\frac{E}{\sqrt{2}} = \frac{E}{\sqrt{\left(1 - \frac{\omega^2}{\omega_0^2}\right)^2 + \frac{1}{Q^2} \frac{\omega^2}{\omega_0^2}}}$$

L'équation bicarrée obtenue admet pour seule solution réelle positive :

$$\omega_1 = \omega_0 \sqrt{1 - \frac{1}{2Q^2} + \sqrt{1 + \left(\frac{1}{2Q^2} - 1\right)^2}}$$

qui définit également la largeur de la bande passante.

- Si $Q > \frac{\sqrt{2}}{2}$: on a alors une valeur maximale $u_{C,\max} = \frac{E}{\sqrt{\frac{1}{Q^2} - \frac{1}{4Q^4}}}$. Ici deux cas

peuvent se présenter :

- $u_{C,\max}$ prend une valeur supérieure à 3 dB : dans ce cas, la bande passante correspond à la largeur de la résonance.
- $u_{C,\max}$ prend une valeur inférieure à 3 dB : on retrouve une bande passante du type $\omega < \omega'_1$.

Pour déterminer plus précisément la bande passante, on pose $x = \frac{\omega}{\omega_0}$ et on résout :

$$u_C = \frac{E}{\sqrt{\frac{x^2}{Q^2} + (1 - x^2)^2}} \geq \frac{u_{C,\max}}{\sqrt{2}}$$

soit en prenant l'inverse et en développant :

$$x^4 + \left(\frac{1}{Q^2} - 2\right)x^2 + \left(1 - \frac{2}{Q^2} + \frac{1}{2Q^4}\right) \leq 0$$

dont le discriminant $\Delta = \frac{1}{Q^4} (4Q^2 - 1)$ est positif car $Q > \frac{\sqrt{2}}{2}$. On en déduit :

$$x^2 = \frac{2Q^2 - 1 \pm \sqrt{4Q^2 - 1}}{2Q^2}$$

Or $2Q^2 - 1 > \sqrt{4Q^2 - 1}$ pour $Q > \frac{\sqrt{2+\sqrt{2}}}{2}$. On aura donc :

- ◊ Une seule solution x^2 positive pour $\frac{\sqrt{2}}{2} < Q < \frac{\sqrt{2+\sqrt{2}}}{2}$, cela correspond au cas où $u_{C,\max}$ prend une valeur inférieure à 3 dB. La bande passante vaut alors :

$$\Delta\omega = \omega'_1 = \omega_0 \sqrt{1 - \frac{1}{2Q^2} + \frac{1}{Q} \sqrt{1 - \frac{1}{4Q^2}}}$$

- ◊ Deux solutions x^2 positives pour $Q > \frac{\sqrt{2+\sqrt{2}}}{2}$, cela correspond au cas où $u_{C,\max}$ prend une valeur supérieure à 3 dB. Les limites de la bande passante sont :

$$\omega'_1 = \omega_0 \left(\sqrt{1 - \frac{1}{2Q^2} + \frac{1}{Q} \sqrt{1 - \frac{1}{4Q^2}}} \right)$$

et

$$\omega_1 = \omega_0 \sqrt{1 - \frac{1}{2Q^2} - \frac{1}{Q} \sqrt{1 - \frac{1}{4Q^2}}}$$

La largeur de la bande passante est alors :

$$\Delta\omega = \omega'_1 - \omega_1 = \omega_0 \left(\sqrt{1 - \frac{1}{2Q^2}} + \frac{1}{Q} \sqrt{1 - \frac{1}{4Q^2}} - \sqrt{1 - \frac{1}{2Q^2}} - \frac{1}{Q} \sqrt{1 - \frac{1}{4Q^2}} \right)$$

On notera que la bande passante de la résonance aux bornes du condensateur n'est pas la même que celle de la résonance en intensité.

3.4 Cas où $Q \gg 1$ (amortissement faible)

Dans ce cas, on a résonance aux bornes du condensateur : elle a lieu pour $\omega_r = \omega_0 \sqrt{1 - \frac{1}{2Q^2}} \simeq \omega_0$. La fréquence de résonance est la même, que ce soit aux bornes du condensateur ou en intensité.

On étudie

$$u_C = |\underline{u}_C| = \frac{i}{C\omega}$$

au voisinage de la résonance :

$$u_C \simeq \frac{i}{C\omega_0}$$

Or, lors de l'analyse de la résonance en intensité, la bande passante correspondante a été obtenue pour :

$$i > \frac{i_{\max}}{\sqrt{2}} \quad \text{soit} \quad u_C \simeq \frac{i}{C\omega_0} > \frac{i_{\max}}{\sqrt{2}C\omega_0} = \frac{u_{C,\max}}{\sqrt{2}}$$

On en déduit qu'une bonne approximation de la bande passante de la résonance aux bornes du condensateur pour de très bonnes valeurs du facteur de qualité ($Q \gg 1$) est fournie par celle de la résonance en intensité. On en déduit :

$$\Delta\omega \simeq \frac{\omega_0}{Q}$$

On peut également obtenir ce résultat en faisant un développement limité de ω_1 et ω'_1 au premier ordre en $\frac{1}{Q}$:

$$\omega_1 \simeq \omega_0 \sqrt{1 - \frac{1}{2Q^2}} \simeq \omega_0 \left(1 - \frac{1}{2Q} \right) \quad \text{et} \quad \omega'_1 \simeq \omega_0 \sqrt{1 + \frac{1}{Q^2}} \simeq \omega_0 \left(1 + \frac{1}{2Q} \right)$$

dont on déduit une expression approchée de la bande passante :

$$\Delta\omega = \omega'_1 - \omega_1 \simeq \omega_0 \left(1 + \frac{1}{2Q} - 1 + \frac{1}{2Q} \right) = \frac{\omega_0}{Q}$$

Dans le cas de bons facteurs de qualité, la pulsation de résonance en intensité et aux bornes du condensateur est sensiblement la même et la tension aux bornes du condensateur est la valeur maximale. Si $Q \gg 1$,

$$U_{C,\max} \simeq QE \quad \text{ou encore} \quad Q \simeq \frac{U_{C,\max}}{E}$$



Il est important de bien remarquer que l'expression de Q en fonction de la tension de résonance aux bornes du condensateur n'est valable que pour $Q \gg 1$ alors que le facteur de surtension défini à partir de la tension aux bornes du condensateur pour la pulsation de résonance **en intensité** est valable pour toute valeur du facteur de qualité.

L'étude de la résonance en tension aux bornes de la bobine s'effectue de manière analogue à ce qui vient d'être fait aux bornes du condensateur.

4. Étude de l'impédance

On peut compléter l'analyse de la résonance du circuit R, L, C série par l'étude de l'impédance du circuit.

Quelle que soit la tension étudiée et donc le montage pratique employé, l'impédance totale du circuit est la même à savoir celle de l'association en série de R , de L et de C soit :

$$\underline{Z} = R + jL\omega + \frac{1}{jC\omega} = R + j \left(L\omega - \frac{1}{C\omega} \right)$$

4.1 Impédances limites

Avant de détailler l'étude par les calculs, on peut obtenir les comportements dits asymptotiques de l'impédance en précisant ce qui se passe à hautes et basses fréquences ainsi qu'à la fréquence particulière obtenue lors de l'analyse de la résonance : la fréquence propre $f_0 = \frac{\omega_0}{2\pi}$.

On rappelle les comportements en basses et hautes fréquences d'une bobine et d'un condensateur :

- en basses fréquences, le condensateur se comporte comme un interrupteur ouvert et la bobine comme un fil ;
- en hautes fréquences, c'est l'inverse : le condensateur se comporte comme un fil et la bobine comme un interrupteur ouvert.

Cela permet d'en déduire les comportements limites de l'impédance en fonction de la pulsation :

- quand ω tend vers 0, le terme prépondérant dans l'expression de l'impédance est $\frac{1}{jC\omega}$: le circuit aura donc à basses fréquences un comportement de type capacitif,

- quand ω tend vers $\omega_0 = \frac{1}{\sqrt{LC}}$, le terme prépondérant dans l'expression de l'impédance est R : le circuit aura donc un comportement de type résistif,
- quand ω tend vers $+\infty$, le terme prépondérant dans l'expression de l'impédance est $jL\omega$: le circuit aura donc à hautes fréquences un comportement de type inductif.

Le comportement de l'association en série d'une résistance, d'une inductance et d'une capacité évolue donc avec la fréquence.

4.2 Étude du module

Le module s'exprime par :

$$|\underline{Z}| = \sqrt{R^2 + \left(L\omega - \frac{1}{C\omega}\right)^2}$$

On cherche à étudier les variations de cette quantité avec la fréquence ou avec la pulsation ω . On remarque qu'il s'agit de l'inverse de l'amplitude de l'intensité au facteur E près. Par conséquent, les variations du module sont opposées à celle de l'amplitude de l'intensité à savoir le module de l'impédance passe par un minimum égal à R pour la pulsation propre ou pulsation de résonance.

On peut remarquer que :

- quand ω tend vers 0, $|\underline{Z}|$ tend vers $+\infty$ comme $\frac{1}{C\omega}$ qui est le terme prépondérant ;
- quand ω tend vers $+\infty$, $|\underline{Z}|$ tend vers $+\infty$ comme $L\omega$ qui est le terme prépondérant. On note l'existence d'une asymptote pour ω tendant vers l'infini : la droite passant par l'origine et de pente L ;
- quand $\omega = \omega_0 = \frac{1}{\sqrt{LC}}$, $|\underline{Z}| = R$.

Cette étude permet de donner l'allure de $|\underline{Z}|$ en fonction de ω :

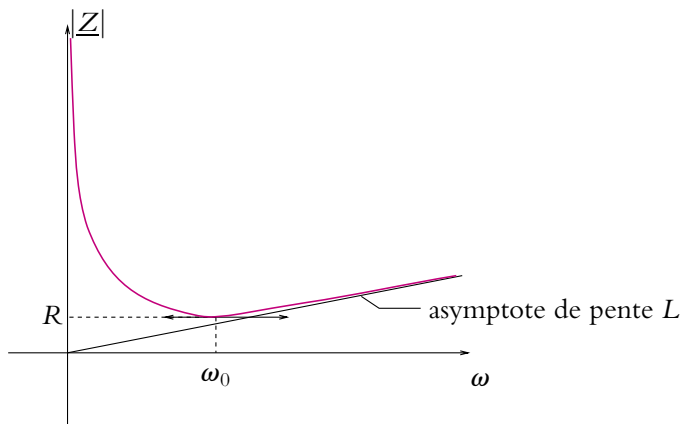


Figure 16.7 Module de l'impédance d'un circuit R, L, C série.

On peut écrire l'impédance en fonction de la pulsation propre et du facteur de qualité $Q = \frac{L\omega_0}{R}$:

$$\underline{Z} = R \left(1 + jQ \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right) \right)$$

4.3 Étude de l'argument

On obtient l'argument par l'expression de sa tangente et de l'intervalle auquel il appartient :

$$\begin{cases} \tan(\text{Arg}(\underline{Z})) = Q \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right) \\ \cos \varphi > 0 \end{cases} \quad \text{donc} \quad \begin{cases} \tan(\text{Arg}(\underline{Z})) = Q \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right) \\ \text{Arg}(\underline{Z}) \in \left[-\frac{\pi}{2}, \frac{\pi}{2} \right] \end{cases}$$

On remarque que l'argument de l'impédance est l'opposé du déphasage entre l'intensité et la tension imposée.

On en déduit donc que l'argument de $\underline{Z}$ est une fonction croissante de ω .

D'autre part, quand ω tend vers 0, $\tan(\text{Arg}(\underline{Z}))$ tend vers $-\infty$ avec $\text{Arg}(\underline{Z}) \in \left[-\frac{\pi}{2}, \frac{\pi}{2} \right]$, donc $\text{Arg}(\underline{Z})$ tend vers $-\frac{\pi}{2}$.

Quand ω tend vers $+\infty$, $\tan(\text{Arg}(\underline{Z}))$ tend vers $+\infty$ avec $\text{Arg}(\underline{Z}) \in \left[-\frac{\pi}{2}, \frac{\pi}{2} \right]$, donc $\text{Arg}(\underline{Z})$ tend vers $+\frac{\pi}{2}$.

Quand $\omega = \omega_0$, $\tan(\text{Arg}(\underline{Z}))$ tend vers 0 avec $\text{Arg}(\underline{Z}) \in \left[-\frac{\pi}{2}, \frac{\pi}{2} \right]$, donc $\text{Arg}(\underline{Z})$ tend vers 0.

Cette étude permet de donner l'allure de $\text{Arg}(\underline{Z})$ en fonction de ω (voir figure 16.8).

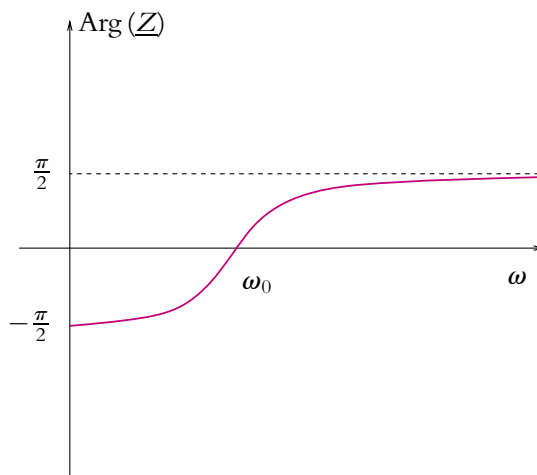


Figure 16.8 Argument de l'impédance d'un circuit *RLC* série.

A. Applications directes du cours

1. Résonance d'un circuit R, L, C parallèle

On considère le circuit suivant :

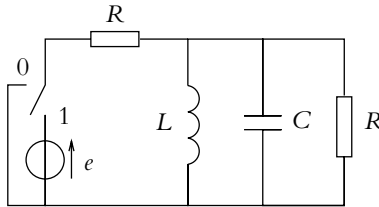


Figure 16.9

où e est une tension sinusoïdale de pulsation ω .

- Donner l'expression complexe de la tension s aux bornes de l'association en parallèle R, L, C .
- Etablir qu'il y a un phénomène de résonance pour la tension s . On précisera la pulsation à laquelle ce phénomène se produit.
- Déterminer la bande passante correspondante.
- En déduire l'expression du facteur de qualité.
- Que peut-on dire du déphasage à la résonance de la tension s ?
- Comparer cette résonance avec la résonance en intensité d'un circuit R, L, C série.

2. Recherche d'un circuit résonant

On dispose d'une inductance $L = 40,0 \text{ mH}$, d'une résistance R variable entre $1,00 \text{ k}\Omega$ et $100 \text{ k}\Omega$ d'une capacité variable entre $10,0 \text{ nF}$ et $1,00 \mu\text{F}$.

On veut réaliser un circuit résonant à la fréquence $f_0 = 2,00 \text{ kHz}$ de gain 1 à cette fréquence.

- Proposer des réalisations possibles d'un tel montage.
- Étudier la sélectivité de chacun d'eux.
- Préciser les valeurs de R et C à choisir pour obtenir la bonne fréquence de résonance et la meilleure sélectivité de résonance.
- Quelle est la "meilleure" solution quant à la sélectivité ?

B. Exercices et problèmes

1. Étude d'une résonance (1)

On considère le circuit ci-dessous alimenté par une source de tension sinusoïdale de f.e.m. $e(t) = E\sqrt{2} \cos(\omega t)$.

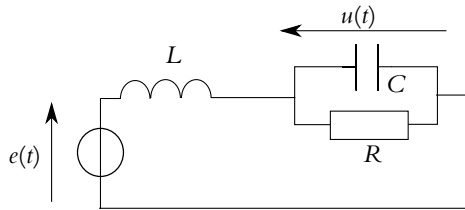


Figure 16.10

1. Obtention de l'expression de $u(t)$ par deux méthodes

- Première méthode* : en utilisant un diviseur de tension, établir l'expression complexe de $\underline{u}$ en fonction de E , L , R , C et ω .
- Deuxième méthode* : en transformant le générateur de Thévenin (e , L) en modèle de Norton retrouver l'expression de $\underline{u}$.

2. Étude de la réponse en fréquence du circuit

Étudier le module de $\underline{u}$ et son argument en fonction de ω . Tracer les courbes.

2. Étude d'une résonance (3)

Soit le circuit suivant :

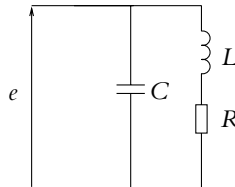


Figure 16.11

où e est une tension sinusoïdale de pulsation ω . On étudie l'intensité parcourant le générateur.

- Exprimer l'impédance complexe du circuit.

- On pose $\omega_0 = \frac{1}{\sqrt{LC}}$, $Q = \frac{L\omega_0}{R}$ et $x = \frac{\omega}{\omega_0}$.

Donner l'expression de l'impédance en fonction de x , Q et R .

- Établir l'existence d'un extremum du module de l'impédance pour certaines valeurs de Q qu'on précisera.
- Donner l'expression de la pulsation correspondant à l'extremum.
- En étudiant les limites du module de l'impédance, en déduire qu'il s'agit d'un maximum.
- Que peut-on en déduire pour l'intensité parcourant le générateur ?
- Dans le cas où Q est grand, donner les expressions approchées de la pulsation de résonance en impédance et de la valeur correspondante du maximum de $|\underline{Z}|$.

3. Quelques résonances, d'après Centrale TSI 2003

1. Résonance série :

- Écrire l'impédance $\underline{Z}'$ d'un circuit composé d'une résistance R , d'une inductance L et d'une capacité C placées en série.
- Exprimer le module Z' de l'impédance en fonction de $\omega_0 = \frac{1}{\sqrt{LC}}$ et de $Q = \frac{L\omega_0}{R}$.
- Déterminer le retard de phase φ de l'intensité i parcourant ce circuit lorsqu'on l'alimente par une tension sinusoïdale d'amplitude E et de pulsation ω .
- Tracer en le justifiant le rapport $\frac{Z'}{R}$ en fonction de $\frac{\omega}{\omega_0}$.
- En déduire que le module de l'intensité I passe par un maximum $I_{\max}$ dont on donnera la valeur.
- Tracer le rapport $\frac{I}{I_{\max}}$ en fonction de $\frac{\omega}{\omega_0}$.
- Même question pour φ .
- On définit l'acuité de la résonance par $A = \frac{\Delta\omega}{\omega_0}$ où $\Delta\omega = \omega_2 - \omega_1$ est la bande de pulsations pour laquelle $I(\omega) \geq \frac{I_{\max}}{\sqrt{2}}$. Exprimer A en fonction de Q .
- En déduire une interprétation physique des paramètres ω_0 et Q .
- Dans quel domaine varie la phase φ pour une pulsation appartenant à la bande de pulsation $\Delta\omega$?

2. Résonance parallèle :

- Écrire l'impédance $\underline{Z}$ du nouveau circuit composé de l'association en parallèle d'un condensateur de capacité C et d'une bobine d'inductance L et de résistance R .
- Exprimer $\underline{Z}$ en fonction de Q , ω , ω_0 , $\underline{Z}'$, C et R .
- Montrer que pour $Q \gg 1$, on peut écrire $\underline{Z} = \frac{R^2 Q^2}{\underline{Z}'}$.
- Quelle est la valeur de $\underline{Z}$ lorsque la pulsation tend vers ω_0 ?
- Que peut-on en conclure?
- Pour $\omega = \omega_0$, donner l'expression du courant i_1 traversant le condensateur.
- Même question pour le courant i_2 traversant la bobine.

4. Analyse expérimentale d'un circuit R, L, C série, d'après CCP DEUG 2005

Soit un circuit R, L, C série alimenté par une tension sinusoïdale dont on regarde la tension e aux bornes du générateur en voie 2 et s aux bornes de R en voie 1.

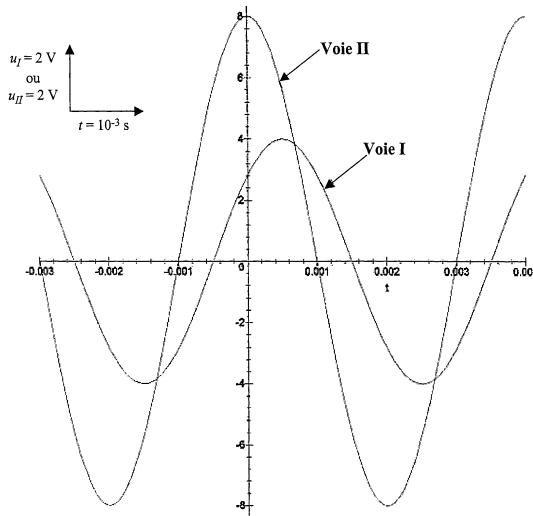


Figure 16.12

1. Sélectivité d'un filtre passe-bande :

- Établir l'expression de l'impédance complexe $\underline{Z}$ du circuit R, L, C série.
- Exprimer la fonction de transfert du montage en fonction de R, L, C et ω .
- On pose $\omega_0 = \frac{1}{\sqrt{LC}}$, $x = \frac{\omega}{\omega_0}$ et $Q = \frac{L\omega_0}{R}$. Donner l'expression de la fonction de transfert en fonction de ces nouveaux paramètres.
- Montrer que, pour toute valeur de Q , le gain admet un même maximum $G_{\max}$. On précisera la pulsation correspondante.
- Déterminer la pulsation pour laquelle le déphasage est nul.
- Rappeler la définition de la bande passante et déterminer la largeur de celle-ci pour le montage considéré.
- En déduire que la sélectivité du filtre est liée au choix de R pour L et C fixées.
- Tracer les courbes donnant le gain en fonction de ω pour deux valeurs Q_1 et Q_2 du facteur de qualité avec $Q_1 > Q_2$ en faisant apparaître les bandes passantes.

2. Détermination des paramètres de la bobine :

Pour $R = 20 \, \Omega$ et $C = 10 \, \mu\text{F}$, on a les oscillogrammes donnés par la figure 16.12.

- Déterminer graphiquement les valeurs de la période T , de la pulsation ω , de l'amplitude de l'intensité I_M , de l'amplitude de la tension délivrée par le générateur U_M et de la norme de l'impédance Z du circuit.
- Quelle voie est en retard sur l'autre ?
- Déterminer le déphasage φ entre les deux tensions.

- d) Montrer que les valeurs expérimentales de Z , φ et R sont incohérentes dans l'hypothèse d'une bobine idéale modélisée par une seule inductance L .
- e) On prend comme modèle de la bobine une inductance L en série avec une résistance r . Déterminer la valeur de r à partir des observations expérimentales.
- f) En déduire la valeur de L .

5. Étude du modèle simplifié d'un moteur, d'après Banque PT 2004

Un moteur est alimenté par une tension sinusoïdale d'amplitude complexe $\underline{V}$; on note $\underline{I}$ l'amplitude complexe de l'intensité du courant passant dans le circuit.

Le schéma électrique équivalent est représenté sur la figure (16.13). La source de tension alimentant le moteur est supposée parfaite.

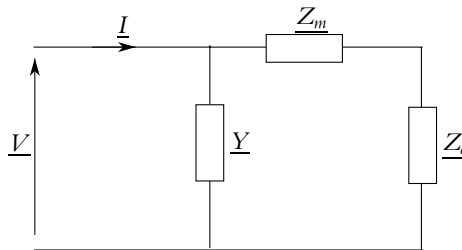


Figure 16.13

$\underline{Y}$ représente l'admittance électrique complexe intrinsèque du moteur. On admettra que le comportement mécanique du moteur se traduit du point de vue électrique par une impédance complexe $\underline{Z}_m$ (dite impédance motionnelle) et par $\underline{Z}_c$, dont la valeur est fonction de la charge mécanique du moteur.

Le dipôle d'admittance $\underline{Y}$ est constitué d'une résistance R_0 en parallèle avec un condensateur C_0 . L'impédance $\underline{Z}_m$ correspond à un circuit RLC série. L'impédance $\underline{Z}_c$ correspond à une résistance R_c .

1. Représenter la figure (16.13) en remplaçant les éléments $\underline{Y}$, $\underline{Z}_m$ et $\underline{Z}_c$ par les inductances, résistances et condensateurs qui leur correspondent.
2. Déterminer la pulsation de résonance en courant ω_s du circuit série constitué de $\underline{Z}_m$ et $\underline{Z}_c$.

Pour la suite, on prendra les valeurs numériques $R_0 = 18\text{k}\Omega$, $C_0 = 8\text{ nF}$, $R = 50\ \Omega$, $L = 0,1\text{ H}$, $C = 0,2\text{ nF}$ et $R_c = 50\ \Omega$.

3. On note $\underline{Y}_p$ l'admittance complexe du circuit constitué de l'ensemble des éléments $\underline{Z}_m$, $\underline{Z}_c$ et $\underline{Y}$. Les figures (16.14) et (16.15) représentent l'évolution du module Y_p de l'admittance complexe $\underline{Y}_p$ en fonction de la pulsation réduite $x = \frac{\omega}{\omega_s}$. La figure (16.15) est un agrandissement de la figure (16.14) autour de $x = 1$

À partir des figures (16.14) et (16.15), déterminer la valeur numérique de la fréquence (en hertz) qui permet d'obtenir une résonance en courant.

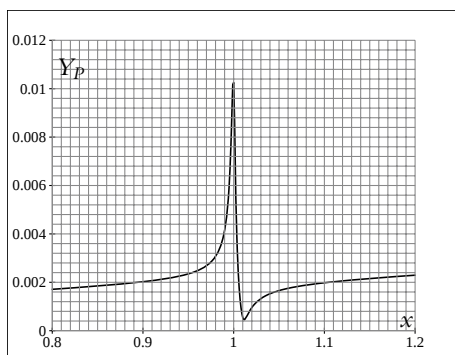


Figure 16.14

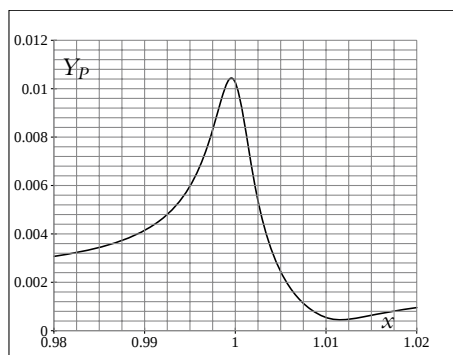


Figure 16.15

4. Comparer $Y = |Y|$ et le module de l'admittance complexe $\underline{Y}' = \frac{1}{\underline{Z}_m + \underline{Z}_c}$ lorsque $\omega = \omega_s$ et interpréter le résultat de la question précédente.
5. Une modification du poids de la charge transportée par le moteur provoque une variation de la résistance R_c , qui reste toutefois de l'ordre de grandeur de la dizaine d'ohms ; cette modification a-t-elle un effet significatif sur la valeur de la résonance en courant ? Justifier la réponse.

17

Puissance

1. Définitions

La notion de puissance instantanée a été introduite dans le chapitre sur les circuits linéaires. Le cas particulier des circuits alimentés en régime sinusoïdal est important, car les appareils de la vie courante sont alimentés par la tension sinusoïdale du réseau électrique. Il est nécessaire tout d'abord de donner ou rappeler quelques définitions. La puissance est aussi à la base de la définition de la valeur efficace d'une intensité ou d'une tension.

1.1 Puissance instantanée

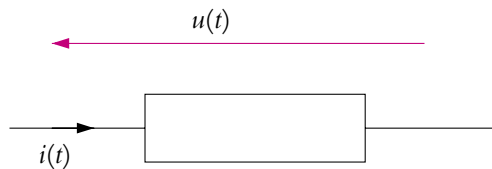


Figure 17.1 Définition de la puissance en convention récepteur.

On considère un dipôle en convention récepteur (Cf. figure 17.1). La puissance instantanée reçue par le dipôle est :

$$\mathcal{P}(t) = u(t) i(t)$$

Pendant un intervalle de temps dt , le dipôle reçoit une petite quantité d'énergie : $\delta\mathcal{E} = \mathcal{P}dt$. Entre les instants t_1 et t_2 , il reçoit l'énergie :

$$\mathcal{E} = \int_{t_1}^{t_2} \delta\mathcal{E} = \int_{t_1}^{t_2} \mathcal{P}(t)dt$$

Si le dipôle est en convention générateur (Cf. figure 17.2), on parlera alors de puissance cédée par le dipôle :

$$\mathcal{P}(t) = u'(t)i(t)$$

On peut, comme pour la puissance reçue, calculer l'énergie cédée par le dipôle.

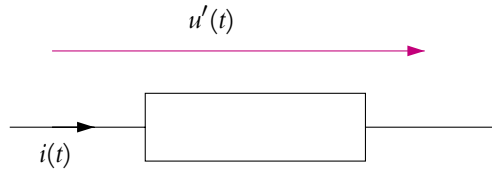


Figure 17.2 Définition de la puissance en convention générateur.

1.2 Puissance moyenne

En général, ce n'est pas la puissance à un instant donné qui intéresse le consommateur. Si on prend l'exemple de la puissance consommée dans une habitation, sachant que le secteur délivre un courant à une fréquence de 50 Hz, on ne peut pas mesurer la puissance tous les $1/50^{\text{e}}$ de seconde. On s'intéresse plutôt à la puissance consommée sur une longue période : il s'agit de la *puissance moyenne*. Ainsi la puissance moyenne reçue par un dipôle sur un intervalle de temps τ est :

$$\mathcal{P}_m = \frac{1}{\tau} \int_0^{\tau} \mathcal{P}(t) dt \quad (17.1)$$

Dans la plupart des cas rencontrés, la tension et l'intensité sont des fonctions périodiques du temps de période T et même souvent sinusoïdales. Dans ce cas, il suffit d'effectuer uniquement l'intégrale précédente sur une période T du phénomène :

$$\mathcal{P}_m = \frac{1}{T} \int_0^T \mathcal{P}(t) dt \quad (17.2)$$

1.3 Valeurs moyenne et efficace d'une tension ou d'une intensité

a) Valeur moyenne

Comme pour une puissance, on désire souvent connaître la valeur moyenne d'une tension ou d'une intensité. Dans la suite, les expressions seront établies pour une tension $u(t)$, mais les définitions et les résultats sont semblables pour une intensité $i(t)$. Si l'on se restreint à un signal périodique de période T , la valeur moyenne du signal $u(t)$ (notée $\bar{u}$ ou $\langle u \rangle$) sera définie par :

$$\bar{u} = \langle u \rangle = \frac{1}{T} \int_0^T u(t) dt \quad (17.3)$$

La différence entre $u(t)$ et $\bar{u}$ est notée $u_{\text{ond}}(t)$, l'indice « ond » signifiant *ondulation*, soit :

$$u(t) = \bar{u} + u_{\text{ond}}(t)$$

Que représentent physiquement $\bar{u}$ et u_{ond} ? Si l'on se réfère à la figure 17.3, la valeur moyenne représente une valeur continue (en pointillés sur la figure) autour de laquelle oscille le signal $u(t)$. Cette oscillation porte le nom d'ondulation ou partie alternative du signal.

Un signal périodique quelconque est donc la somme d'un signal continu (on parle de tension de décalage ou d'offset) et d'une ondulation.

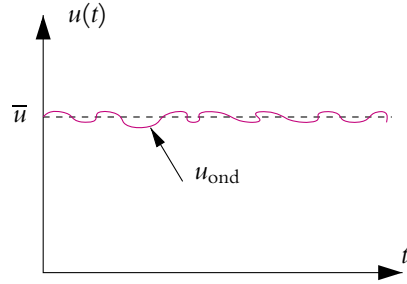


Figure 17.3 Valeur moyenne et ondulation d'un signal.

On peut donner trois exemples de valeurs moyennes (signaux de la figure 17.4) :

- un signal continu : $u(t) = U_c$;
- un signal sinusoïdal pur : $u(t) = U \cos(\omega t)$;
- un signal créneau :
 - $u(t) = 2U$ si $nT < t < \left(n + \frac{1}{2}\right) T$
 - $u(t) = U$ si $\left(n + \frac{1}{2}\right) T < t < (n + 1) T$.

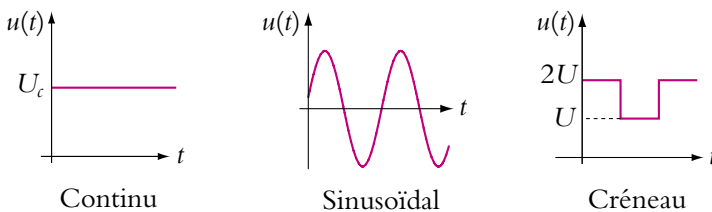


Figure 17.4 Valeurs moyennes de quelques signaux.

1. Pour le signal continu, $\bar{u} = U_c$ et $u_{\text{ond}} = 0$.
2. Pour le signal sinusoïdal pur :

$$\bar{u} = \langle u \rangle = \frac{1}{T} \int_0^T U \cos(\omega t) dt = 0$$

et $u_{\text{ond}} = u(t)$.

3. Pour le signal créneau :

$$\bar{u} = \langle u \rangle = \frac{1}{T} \left(\int_0^{T/2} 2U dt + \int_{T/2}^T U dt \right) = \frac{3U}{2}$$

et :

$$u_{\text{ond}}(t) = \frac{U}{2} \quad \text{si} \quad nT \leq t \leq \left(n + \frac{1}{2}\right) T$$

$$u_{\text{ond}}(t) = -\frac{U}{2} \quad \text{si} \quad \left(n + \frac{1}{2}\right) T \leq t \leq (n+1)T$$

b) Valeur efficace :

La définition de la *valeur efficace* d'un signal (tension ou intensité) est d'origine énergétique :

La valeur efficace d'une tension $u(t)$ (respectivement d'une intensité $i(t)$) est la valeur de la tension continue (respectivement intensité continue) qui recevrait la même énergie pendant la même durée dans le même résistor.

Ainsi sur l'intervalle de temps τ , l'énergie reçue par le résistor de résistance R est d'une part :

$$\mathcal{E}_{0 \rightarrow \tau} = \int_0^{\tau} \frac{u(t)^2}{R} dt$$

et d'autre part, si on note U_{eff} la valeur efficace :

$$\mathcal{E}_{0 \rightarrow \tau} = \int_0^{\tau} \frac{U_{\text{eff}}^2}{R} dt = \frac{U_{\text{eff}}^2}{R} \tau$$

En égalant les deux expressions, on obtient celle de la valeur efficace :

$$U_{\text{eff}} = \sqrt{\frac{1}{\tau} \int_0^{\tau} u(t)^2 dt} \quad (17.4)$$

On s'intéressera à des signaux périodiques et, comme cela a été signalé auparavant, il suffit de calculer l'intégrale précédente sur une période T du signal. Dans ce cas, on obtient :

$$U_{\text{eff}} = \sqrt{\frac{1}{T} \int_0^T u(t)^2 dt} \quad (17.5)$$

En comparant avec l'équation (17.3), on constate que la valeur efficace n'est autre que **la racine carrée de la valeur moyenne du carré**, soit *Root Mean Square* en anglais. On parlera donc de valeur efficace ou de valeur R.M.S. Le chapitre sur l'instrumentation électrique détaillera ce que mesurent effectivement les appareils.

On peut calculer les valeurs efficaces des signaux de la figure 17.4 :

1. pour le signal continu : $U_{\text{eff}} = U_{\text{RMS}} = U_c$;
2. pour le signal sinusoïdal :

$$U_{\text{eff}} = U_{\text{RMS}} = \sqrt{\frac{1}{T} \int_0^T U^2 \cos^2(\omega t) dt} = \frac{U}{\sqrt{2}}$$

(ce dernier résultat est bien connu) ;

3. pour le créneau :

$$U_{\text{eff}} = U_{\text{RMS}} = \sqrt{\frac{1}{T} \left(\int_0^{T/2} 4U^2 dt + \int_{T/2}^T U^2 dt \right)} = U\sqrt{\frac{5}{2}}$$

On remarque que le facteur $1/\sqrt{2}$ ne se généralise pas à toutes les formes de signaux.

2. Puissance en régime sinusoïdal forcé

2.1 Puissance instantanée

Les raisonnements seront faits sur la puissance reçue (dipôle en convention récepteur) mais sont généralisables à la puissance cédée. Soit le dipôle de la figure 17.1. On note $u(t)$ la tension à ses bornes et $i(t)$ le courant le parcourant (en convention récepteur) : $u(t) = U\sqrt{2} \cos(\omega t)$ et $i(t) = I\sqrt{2} \cos(\omega t - \varphi)$.

Dans les expressions de $u(t)$ et $i(t)$, U et I sont-elles les amplitudes ou les valeurs efficaces ? Puisqu'il y a le facteur $\sqrt{2}$, il s'agit de grandeurs efficaces, et les amplitudes valent $U\sqrt{2}$ et $I\sqrt{2}$.

On utilise des grandeurs réelles et non complexes, ce qui sera justifié ultérieurement.

La puissance instantanée reçue par le dipôle est donc :

$$\mathcal{P}(t) = u(t)i(t) = 2UI \cos(\omega t) \cos(\omega t - \varphi)$$

Soit en utilisant les formules de trigonométrie :

$$\mathcal{P}(t) = UI(\cos(2\omega t - \varphi) + \cos \varphi) \quad (17.6)$$

On remarque que la puissance instantanée est de pulsation 2ω , double de la pulsation de la tension $u(t)$ et de l'intensité $i(t)$. Elle oscille donc deux fois plus rapidement. Ses valeurs extrêmes sont obtenues pour $\cos(2\omega t - \varphi) = \pm 1$ soit :

$$\mathcal{P}_{\text{max}} = UI(1 + \cos \varphi) \quad \text{et} \quad \mathcal{P}_{\text{min}} = UI(-1 + \cos \varphi)$$

Quel est le comportement des trois dipôles R , L et C ?

Dans le cas d'un résistor, le déphasage entre u et i est nul ($\varphi = 0$) donc $\cos \varphi = 1$ donc $0 \leq \mathcal{P} \leq 2UI$.

La puissance reçue est toujours positive et donc le résistor se comporte toujours comme un récepteur.

Dans le cas d'une bobine ($\varphi = \pi/2$) ou d'un condensateur ($\varphi = -\pi/2$), $\cos(\varphi) = 0$, donc $-UI \leq \mathcal{P} \leq UI$. Ainsi, $\mathcal{P}(t) = UI \cos(2\omega t - \varphi)$ et pendant une demi-période la puissance reçue est positive, puis pendant l'autre demi-période elle est négative.

Les bobines et les condensateurs se comportent alternativement en récepteur (ils reçoivent effectivement de l'énergie) et en générateur (ils fournissent effectivement de l'énergie).

2.2 Puissance moyenne ou active

Comme cela a été signalé précédemment, pour connaître la puissance globale reçue (ou consommée) par un dipôle, il est nécessaire de calculer la puissance moyenne. Dans le cas de la puissance instantanée (17.6), on obtient :

$$\mathcal{P}_m = \frac{UI}{T} \left(\int_0^T \cos(2\omega t - \varphi) dt + \int_0^T \cos \varphi dt \right)$$

La première intégrale est nulle (on intègre un cosinus sur deux fois sa période), ce qui donne :

$$\mathcal{P}_m = UI \cos \varphi \quad (17.7)$$

Cette puissance moyenne porte le nom de *puissance active*. Elle est le produit **des valeurs efficaces** de la tension et de l'intensité, et du cosinus du déphasage $\cos \varphi$ appelé *facteur de puissance*.

La puissance active des trois dipôles R , L et C peut être aisément déterminée :

- pour R : $\varphi = 0$ donc $\mathcal{P}_m = UI$, cette valeur est la même qu'en continu et elle est toujours positive ;
- pour L : $\varphi = \pi/2$ donc $\mathcal{P}_m = 0$;
- pour C : $\varphi = -\pi/2$ donc $\mathcal{P}_m = 0$.

Ces trois résultats confirment les conclusions sur les comportements des dipôles du paragraphe précédent.

L'expression (17.7) est souvent d'utilisation peu pratique à cause de la présence de $\cos \varphi$ et il est préférable d'utiliser l'une des expressions démontrées dans la suite.

L'impédance d'un dipôle s'exprime de la manière suivante :

$$\underline{Z} = R(\omega) + jX(\omega) = Z \exp(j\varphi)$$

où $R(\omega)$ et $X(\omega)$ sont la résistance et la réactance du dipôle (Z et φ dépendent aussi de ω , on ne le précise pas pour ne pas alourdir les notations). Ainsi on peut exprimer $\cos \varphi$ par :

$$\cos \varphi = \frac{R(\omega)}{Z}$$

D'autre part, les valeurs efficaces U et I sont reliées par $U = ZI$. Ainsi en combinant les équations précédentes, on obtient :

$$\mathcal{P}_m = R(\omega)I^2 = \operatorname{Re}(\underline{Z})I^2 \quad (17.8)$$

C'est donc la partie réelle de l'impédance qui consomme l'énergie ce qui confirme qu'une bobine et un condensateur n'en consomment pas en moyenne. On peut établir une expression faisant intervenir la conductance. Le raisonnement est similaire. L'admittance d'un dipôle s'exprime par :

$$\underline{Y} = G(\omega) + jB(\omega) = Y \exp(-j\varphi)$$

où G et B sont la conductance et la susceptance du dipôle. Ainsi on peut exprimer $\cos \varphi$ par :

$$\cos \varphi = \frac{G(\omega)}{Y}$$

D'autre part, les valeurs efficaces U et I sont reliées par $YU = I$. Ainsi en combinant les équations précédentes, on obtient :

$$\mathcal{P}_m = G(\omega)U^2 = \operatorname{Re}(\underline{Y})U^2 \quad (17.9)$$

Les relations (17.8) et (17.9) sont en général plus faciles à utiliser que (17.7) car il est plus rapide de calculer l'impédance ou l'admittance que le facteur de puissance. La relation (17.8) faisant intervenir la résistance est à réserver aux dipôles en série alors que l'autre est plutôt à utiliser avec les dipôles en parallèle.

Par définition, un dipôle passif ne peut que recevoir effectivement de l'énergie donc sa puissance active est positive ou nulle. D'après les relations (17.8) et (17.9), on en déduit que **les parties réelles de son impédance et de son admittance sont positives ou nulles.**

► **Remarque :** Que se passe-t-il si on utilise les valeurs complexes de la tension et de l'intensité ? Dans ce cas :

$$\underline{u} = U\sqrt{2} \exp(j\omega t) \quad \underline{i} = I\sqrt{2} \exp(j(\omega t - \varphi))$$

soit la valeur moyenne du produit $\underline{u}\underline{i}$:

$$\langle \underline{u}\underline{i} \rangle = \frac{2UI}{T} \exp(-j\varphi) \int_0^T \exp(2j\omega t) dt = \frac{2UI}{2j\omega T} \exp(-j\varphi) [\exp(2j\omega t)]_0^T = 0$$

On obtiendrait donc une puissance moyenne toujours nulle, ce qui est évidemment **faux**.



Le passage aux notations complexes est réservé aux grandeurs linéaires puisque c'est la linéarité des équations qui justifie son utilisation. Or $\mathcal{P}$ est le produit de deux fonctions sinusoïdales, c'est une grandeur non linéaire, **on ne doit donc pas utiliser les notations complexes pour la puissance**¹.

2.3 Calcul du facteur de puissance d'une installation

On considère une installation électrique (Cf. figure 17.5) constituée d'un moteur consommant une puissance active $\mathcal{P}_1$ et de facteur de puissance $\cos \varphi_1 = 0,8$ branché en parallèle sur des lampes dont l'ensemble est modélisé par une résistance $R = 100 \, \Omega$. On note $\mathcal{P}_2$ la puissance active consommée par la résistance R .

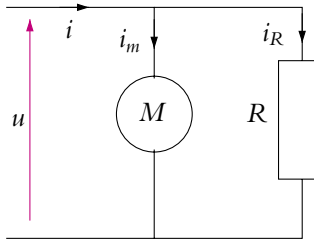


Figure 17.5 Installation électrique.

L'installation est alimentée par le secteur sous une tension efficace U de 240 V. On désire calculer le facteur de puissance de l'ensemble. On applique la loi des nœuds en A :

$$\underline{i} = \underline{i_m} + \underline{i_R}$$

avec :

$$\underline{i} = I\sqrt{2} \exp(j(\omega t - \varphi)), \quad \underline{i_m} = I_m\sqrt{2} \exp(j(\omega t - \varphi_1)) \quad \text{et} \quad \underline{i_R} = I_R\sqrt{2} \exp(j\omega t)$$

¹ Il existe une définition complexe de la puissance, mais elle n'est pas au programme.

► **Remarque** : Dans les expressions précédentes des intensités, puisque le facteur $\sqrt{2}$ apparaît, I , I_m et I_R sont les valeurs efficaces des intensités.

Après simplification des termes dépendant du temps et des $\sqrt{2}$, la loi des nœuds devient :

$$I \exp(-j\varphi) = I_m \exp(-j\varphi_1) + I_R$$

En prenant le module de cette expression, on obtient :

$$I = \sqrt{I_m^2 + I_R^2 + 2I_m I_R \cos \varphi_1}$$

On exprime les différentes intensités en fonction des puissances : $\mathcal{P}_1 = UI_m \cos \varphi_1$, $\mathcal{P} = UI \cos \varphi$ et $\mathcal{P}_2 = UI_R$, d'où :

$$I = \frac{\mathcal{P}}{U \cos \varphi} = \sqrt{\frac{\mathcal{P}_1^2}{U^2 \cos^2 \varphi_1} + \frac{\mathcal{P}_2^2}{U^2} + 2 \cos \varphi_1 \frac{\mathcal{P}_1 \mathcal{P}_2}{U^2 \cos \varphi_1}}$$

soit, après simplification et en utilisant $\frac{1}{\cos^2 \varphi_1} = 1 + \tan^2 \varphi_1$:

$$\cos \varphi = \frac{\mathcal{P}}{\sqrt{(\mathcal{P}_1 + \mathcal{P}_2)^2 + \mathcal{P}_1^2 \tan^2 \varphi_1}} = 0,84 \quad (17.10)$$

2.4 Importance du facteur de puissance

La figure 17.6 schématise la distribution de l'électricité pour une installation alimentée sous une tension efficace U , par une intensité efficace I et consommant une puissance active $\mathcal{P}_m$, avec un facteur de puissance $\cos \varphi$. On appelle R la résistance des câbles amenant le courant jusqu'à l'installation.

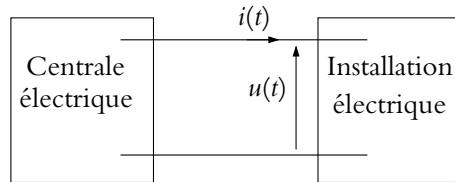


Figure 17.6 Distribution électrique et installation.

Le consommateur paie pour la puissance consommée $\mathcal{P}_m$ mais E.D.F. doit produire une puissance supérieure $\mathcal{P}_p$ en raison de la résistance des fils : $\mathcal{P}_p = \mathcal{P}_m + RI^2$. Le rendement η défini par $\mathcal{P}_m/\mathcal{P}_p$ doit être proche de 1 de manière à ce que la puissance produite soit la plus proche possible de la puissance consommée.

L'intensité efficace alimentant l'installation est :

$$I = \frac{\mathcal{P}_m}{U \cos \varphi}$$

d'où un rendement :

$$\eta = \frac{\mathcal{P}_m}{\mathcal{P}_m + RI^2} = \frac{\mathcal{P}_m}{\mathcal{P}_m + \frac{R\mathcal{P}_m^2}{U^2 \cos^2 \varphi}} \quad (17.11)$$

Pour que η soit proche de 1, il faut que $\frac{RP_m^2}{U^2 \cos^2 \varphi}$ soit le plus petit possible. Sachant que la puissance consommée est fixée, plusieurs solutions sont envisageables :

1. résistance R des câbles petite,
2. tension U élevée,
3. facteur de puissance $\cos \varphi$ le plus élevé possible.

Pour diminuer la résistance des câbles, on choisit de les fabriquer en cuivre qui est le meilleur conducteur (après l'argent, beaucoup plus cher). L'idéal serait de disposer de supraconducteur ($R = 0 \Omega$), mais cela n'existe pas encore à température ambiante. Pour diminuer la résistance des câbles, on pourrait aussi augmenter leur section mais dans ce cas leur masse augmenterait et il faudrait prévoir des pylônes plus solides et ce serait beaucoup plus onéreux.

En ce qui concerne le deuxième point, le transport de l'électricité se fait effectivement à haute tension (380 kV ou 640 kV), puis on transforme la tension à 240 V efficace avant l'installation. Les postes de transformateurs sont d'ailleurs observables en ville en bas des immeubles ou à la campagne sur des poteaux électriques.

Enfin, E.D.F. impose que le facteur de puissance d'une installation soit supérieur à 0,9. Dans le cas contraire, l'utilisateur encourt une pénalité. Ce sont surtout les installations industrielles utilisant de nombreuses machines qui sont en cause car ces machines à moteur électrique ont un comportement plutôt inductif (donc φ proche de $\frac{\pi}{2}$). On va voir dans le paragraphe suivant comment augmenter le facteur de puissance d'une telle installation.

2.5 Amélioration d'un facteur de puissance

On s'intéresse à un moteur électrique de 6 kW, alimenté par le secteur (50 Hz), sous une tension efficace $U = 240$ V, dont le facteur de puissance est $\cos \varphi_m = 0,78$. Il est modélisé par une inductance de $L = 22$ mH et une résistance $R = 5,8 \Omega$ en série. Afin de contrebalancer l'effet inductif du moteur, on cherche le condensateur à placer en parallèle sur le moteur pour que le facteur de puissance de l'ensemble $\cos \varphi$ soit égal à 0,9 ou $\varphi = \pm 0,45$ rad (Cf. figure 17.7).

L'admittance de l'ensemble est :

$$\underline{Y} = jC\omega + \frac{1}{R + jL\omega} = \frac{R}{R^2 + (L\omega)^2} + j \left(C\omega - \frac{L\omega}{R^2 + (L\omega)^2} \right)$$

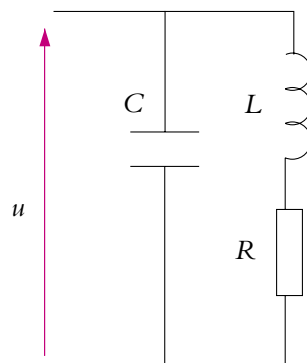


Figure 17.7 Amélioration du facteur de puissance.

Sachant que $\underline{Y} = Y \exp(-j\varphi)$, on tire $\tan(-\varphi)$ des expressions précédentes :

$$\tan(-\varphi) = \frac{C\omega - \frac{L\omega}{R^2 + (L\omega)^2}}{\frac{R}{R^2 + (L\omega)^2}}$$

On tire de l'expression précédente deux valeurs de C qui conviennent : $C = 380 \mu\text{F}$ ou $C = 160 \mu\text{F}$

2.6 Étude énergétique du circuit R, L, C série (PCSI)

a) Équation bilan

On s'intéresse à un circuit R, L, C série connecté à un générateur de f.e.m. e délivrant une puissance $\mathcal{P}_g$. La loi des mailles déjà établie dans les chapitres précédents en régime quelconque conduit à :

$$L \frac{di}{dt} + Ri + u_c = e \quad \text{avec} \quad i = \frac{dq}{dt} \quad \text{et} \quad u_c = \frac{q}{C}$$

En multipliant cette équation par l'intensité i , on obtient :

$$Li \frac{di}{dt} + Ri^2 + \frac{dq}{dt} \frac{q}{C} = ei \quad \Leftrightarrow \quad \frac{d}{dt} \left(\frac{1}{2} Li^2 \right) + \frac{d}{dt} \left(\frac{1}{2} \frac{q^2}{C} \right) + Ri^2 = ei$$

ce qui donne finalement :

$$\frac{d}{dt} \left(\frac{1}{2} Li^2 + \frac{1}{2} \frac{q^2}{C} \right) + Ri^2 = \mathcal{P}_g \quad (17.12)$$

Comme on pouvait s'y attendre, la puissance $\mathcal{P}_g$ fournie par le générateur est égale à la puissance reçue par les trois dipôles R, L et C :

$$\mathcal{P}_L + \mathcal{P}_C + \mathcal{P}_R = \mathcal{P}_g \quad (17.13)$$

$$\text{où } \mathcal{P}_L = \frac{d}{dt} \left(\frac{1}{2} Li^2 \right) \text{ et } \mathcal{P}_C = \frac{d}{dt} \left(\frac{1}{2} \frac{q^2}{C} \right) = \frac{d}{dt} \left(\frac{1}{2} Cu_C^2 \right)$$

b) Puissance consommée en régime sinusoïdal forcé

On s'intéresse maintenant au cas où la f.e.m. est sinusoïdale de la forme $e = E\sqrt{2} \cos \omega t$. Dans ce cas, l'intensité dans le circuit est $i = I\sqrt{2} \cos(\omega t + \varphi)$ et la charge dans le condensateur $q = \frac{I}{\omega} \sqrt{2} \sin(\omega t + \varphi)$ puisque $i = \frac{dq}{dt}$.

Quel est le bilan des puissances moyennes ? Pour cela, il suffit de prendre la valeur moyenne de l'expression (17.13) :

$$< \mathcal{P}_L > + < \mathcal{P}_C > + < \mathcal{P}_R > = < \mathcal{P}_g >$$

Or on a démontré dans les paragraphes précédents que les puissances moyennes reçues par une bobine et un condensateur sont nulles donc :

$$\langle \mathcal{P}_R \rangle = \langle \mathcal{P}_g \rangle \quad (17.14)$$

En régime sinusoïdal forcé, la puissance moyenne fournie par le générateur est consommée dans le résistor par effet Joule.

On s'intéresse maintenant au cas particulier où la pulsation ω est égale à la pulsation propre du circuit $\omega_0 = \frac{1}{\sqrt{LC}}$.

L'énergie instantanée contenue dans L et C est alors :

$$\mathcal{E}_L + \mathcal{E}_C = \frac{1}{2} Li^2 + \frac{1}{2} \frac{q^2}{C} = LI^2 \cos^2(\omega_0 t + \varphi) + \frac{I^2}{C\omega_0^2} \sin^2(\omega_0 t + \varphi)$$

soit avec $\omega_0 = \frac{1}{\sqrt{LC}}$:

$$\mathcal{E}_L + \mathcal{E}_C = LI^2 (\cos^2(\omega_0 t + \varphi) + \sin^2(\omega_0 t + \varphi)) = LI^2$$

L'énergie instantanée reçue par la bobine et le condensateur lorsque $\omega = \omega_0$ est constante donc la puissance instantanée reçue par ces deux dipôles est nulle. Ainsi, lorsque la pulsation est égale à ω_0 , la puissance instantanée fournie par le générateur est consommée uniquement par le résistor. Ce n'est plus seulement vérifié pour les valeurs moyennes.

c) Résonance en puissance

Dans le paragraphe précédent, il a été établi que la puissance moyenne $\mathcal{P}_g$ fournie par le générateur est égale à la puissance moyenne reçue par le résistor soit :

$$\langle \mathcal{P}_g \rangle = \langle \mathcal{P}_R \rangle = RI^2$$

Or, dans le cas d'un circuit R, L, C série, la valeur efficace de l'intensité est :

$$I(\omega) = \frac{E}{R} \frac{1}{\sqrt{1 + Q^2 \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)^2}}$$

où Q est facteur de qualité du circuit. On peut alors écrire :

$$\langle \mathcal{P}_g \rangle(\omega) = \frac{E^2}{R} \frac{1}{1 + Q^2 \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)^2} \quad (17.15)$$

Ainsi, comme l'intensité, la puissance moyenne fournie par le générateur est maximale pour la pulsation $\omega = \omega_0$. On dit qu'il y a *résonance de puissance* à la pulsation ω_0 . La

puissance moyenne maximale fournie par le générateur est alors $\mathcal{P}_{\max} = \frac{E^2}{R}$ et tout se passe comme s'il n'y avait que le résistor (Cf. paragraphe précédent).

On définit la bande passante comme l'intervalle de pulsations ω vérifiant :

$$\langle \mathcal{P}_g \rangle (\omega) \leq \frac{\mathcal{P}_{\max}}{2}$$

ce qui du point de vue de l'intensité correspond à l'intervalle de pulsations ω vérifiant :

$$I(\omega) \leq \frac{I_{\max}}{\sqrt{2}}$$

en raison de la relation $\langle \mathcal{P}_g \rangle = RI^2$

On retrouve la définition de la bande passante donnée lors l'étude des résonances.

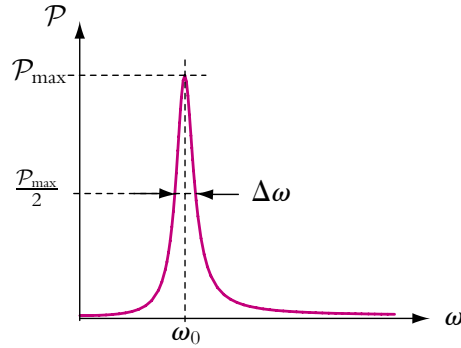


Figure 17.8 Résonance en puissance.

A. Applications directes du cours

1. Détermination d'impédances

Un consommateur reçoit une tension alternative de valeur efficace 240 V. Il utilise cette tension pour alimenter un chauffage électrique qu'on assimilera à une résistance pure et le moteur d'une climatisation qu'on modélisera par une inductance et une résistance en série. Quand seul le chauffage est branché, le compteur débite un courant de 5,0 A. Lorsque seule la climatisation est allumée, le courant est alors de 10 A. Il devient 12 A quand le client branche à la fois son chauffage et sa climatisation. Déterminer les valeurs caractéristiques des appareils.

2. Amélioration du facteur de puissance

Une installation électrique possède les caractéristiques suivantes : $f = 50$ Hz, $I_{\text{eff}} = 20$ A, $U_{\text{eff}} = 380$ V et $P = 5$ kW.

1. Déterminer le facteur de puissance de l'installation.
2. Sachant qu'il s'agit d'un moteur électrique modélisable par une inductance en série avec une résistance, déterminer les valeurs des paramètres.
3. Le fournisseur d'électricité demande à l'utilisateur d'améliorer le facteur de puissance. Pourquoi ?
4. Quel composant l'utilisateur devra employer compte-tenu de la nature de son installation ?
5. Déterminer la valeur de ce composant dans la situation considérée.

3. Mesure de puissance par trois ampèremètres

On souhaite mesurer la puissance d'un appareil. Pour cela, on utilise le montage suivant :

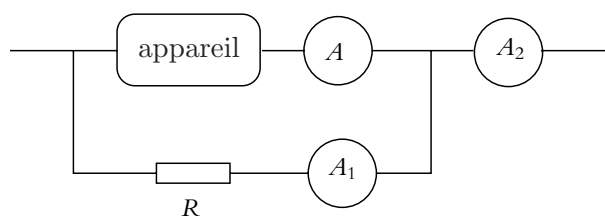


Figure 17.9

Établir l'expression de la puissance de l'appareil en fonction des valeurs lues sur les ampèremètres et de R . On suppose qu'on n'a que des valeurs sinusoïdales centrées.

4. Mesure de puissance par trois voltmètres

Pour mesurer la puissance reçue par un dipôle, on place une résistance R en série avec ce dipôle. On branche trois voltmètres comme sur le schéma.

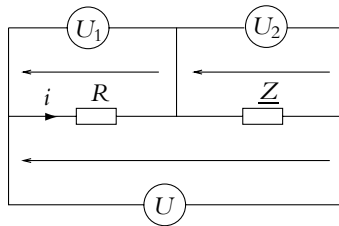


Figure 17.10

Déterminer la puissance reçue par le dipôle en fonction des valeurs efficaces des tensions mesurées et de R .

5. Dipôle L série, CR parallèle, d'après ENAC 2004

Le dipôle AB représenté sur le schéma de la figure (17.11) est alimenté par une source de tension parfaite de force électromotrice instantanée $e(t) = E_0 \sin(\omega t)$.

- Exprimer L en fonction de R , C et ω pour que le dipôle AB soit équivalent à une résistance pure R_{eq} . Exprimer R_{eq} .
- Calculer L sachant que $R = 100 \, \Omega$, $C = 100/3 \, \mu\text{F}$ et $\omega = 400 \, \text{rad.s}^{-1}$.
- Déterminer la valeur efficace I de l'intensité du courant dans la bobine. Calculer sa valeur sachant que la valeur efficace de la force électromotrice du générateur vaut $180 \, \text{V}$.
- Exprimer les valeurs efficaces des différences de potentiel U_{AD} et U_{DB} . Application numérique.
- Déterminer les intensités complexes des courants $\underline{i}_1$ et $\underline{i}_2$ circulant respectivement dans la résistance et dans le condensateur. En déduire leurs valeurs efficaces et leurs déphasages par rapport à la tension entre A et B . Application numérique.
- Calculer la puissance moyenne P sur une période consommée par le dipôle AB .

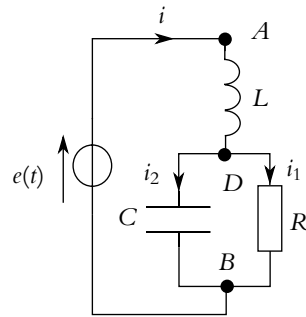


Figure 17.11

B. Exercices et problèmes

1. Facteur de puissance d'une installation

Un moteur consomme une puissance moyenne $P_1 = 4 \, \text{kW}$, son facteur de puissance est $\cos \varphi_1 = 0,8$. Il est branché en parallèle sur une résistance $R = 100 \, \Omega$ modélisant des lampes en parallèle. On note P_2 la puissance active reçue par la résistance. La tension d'alimentation a pour valeur efficace $U = 240 \, \text{V}$, et pour fréquence $50 \, \text{Hz}$.

- Calculer les valeurs efficaces des courants traversant le moteur et la résistance ainsi que la puissance moyenne totale P reçue par l'ensemble.

2. On désire déterminer le facteur de puissance $\cos \varphi$ de l'ensemble.

a) En utilisant la loi des nœuds, établir une expression reliant les valeurs efficaces des intensités. En déduire une expression reliant les puissances et U .

b) Montrer que le facteur de puissance est donné par :

$$\cos \varphi = \frac{P_1 + P_2}{\sqrt{(P_1 + P_2)^2 + P_1^2 \tan^2 \varphi_1}}$$

Calculer la valeur de $\cos \varphi$.

2. Adaptation d'impédance

1. Soit un dipôle d'impédance $\underline{Z}_u$ branché aux bornes d'un générateur de f.e.m. e_g et d'impédance interne $\underline{Z}_g$. Déterminer les conditions (sur les parties réelle et imaginaire de $\underline{Z}_u$) pour que la puissance reçue par l'impédance $\underline{Z}_u$ soit maximale.

2. Soit une installation modélisée par un résistor de résistance R . On souhaite assurer un transfert maximal de puissance du générateur (f.e.m. e_g et résistance interne R_g) vers le dispositif. On fait l'hypothèse que $R > R_g$. Pour cela, on intercale entre le générateur et le dispositif un circuit réalisé avec une inductance L et un condensateur C .

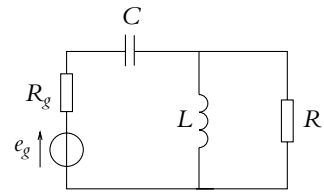


Figure 17.12

Calculer L et C en fonction de R_g , R et de la pulsation ω pour que la puissance reçue par R soit maximale. On dit alors qu'il y a *adaptation d'impédance*.

3. Dans le cas où $R < R_g$, le montage précédent ne convient pas. On se propose de placer L avant C . Calculer de nouveau L et C pour réaliser l'adaptation d'impédance.

3. Adaptation du facteur de puissance d'un moteur, d'après ENAC 2000

Un moteur M équivalent à un résistor de résistance R associé en série avec une bobine de coefficient d'auto-inductance L est alimenté en courant alternatif sinusoïdal de fréquence $f = 50$ Hz par un fil de résistance négligeable (figure 17.13). Le moteur consomme une puissance moyenne $P_m = 4,4$ kW et son facteur de puissance est égal à 0,6. On mesure entre ses bornes A et B une tension de valeur efficace $U = 240$ V.

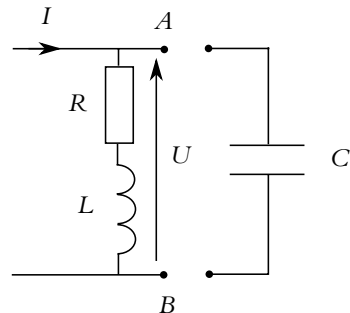


Figure 17.13

1. a) Calculer le courant efficace I circulant dans la ligne.

b) Calculer R et L

2. Pour relever le facteur de puissance de l'installation, on connecte entre les bornes A et B un condensateur de capacité C . La tension mesurée aux bornes du moteur a toujours la valeur efficace $U = 240$ V.

- Calculer la plus petite valeur de C pour que le nouveau facteur de puissance soit égal à 0,9.
- Calculer la puissance moyenne P'_m absorbée par le moteur.
- Calculer l'intensité efficace I' du courant circulant dans la ligne.

4. Rôle d'un fusible

Un dipôle est alimenté en régime sinusoïdal forcé par la tension $u(t) = U\sqrt{2}\cos(2\pi ft)$ avec $U = 240$ V et $f = 50$ Hz. L'intensité du courant qui le parcourt est alors $i(t) = I\sqrt{2}\cos(2\pi ft + \varphi)$.

- Donner les expressions de l'impédance complexe $\underline{Z}$ du dipôle ainsi que celle de son module Z en fonction de U , I et φ (on notera j le nombre imaginaire pur tel que $j^2 = -1$).
- a) Établir l'expression de la puissance électrique moyenne P_e absorbée par le dipôle en fonction de U , I et φ .

- En déduire l'expression $P_e = R \frac{U^2}{Z^2}$, R étant la résistance du dipôle (partie réelle de l'impédance).

Un fusible assimilable à un résistor protège une ligne électrique $(AB, A'B')$, alimentant en régime sinusoïdal forcé, sous une tension efficace U et une fréquence f , un dipôle D assimilable à une bobine d'inductance $L = 30$ mH en série avec un résistor de résistance R (figure 17.14).

L'intensité efficace admissible dans la ligne est $I_{\max} = 16$ A.

Ce dipôle absorbe une puissance électrique moyenne $P_e = 2\,500$ W. La ligne $(AB, A'B')$ se comporte comme un dipôle purement ohmique de résistance électrique totale $R_0 = 1,2\,\Omega$, fusible compris.

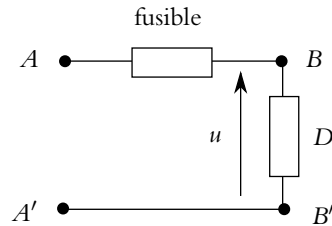


Figure 17.14

- a) Calculer les deux valeurs possibles R_1 et R_2 de la résistance R du dipôle D .
 - b) Pour chaque valeur de R_1 et R_2 , calculer l'intensité efficace I_1 et I_2 dans le dipôle D .
 - c) Déterminer la seule valeur de R possible compte tenu de la présence du fusible.
 - d) En déduire l'intensité efficace I_0 du courant électrique circulant dans la ligne $(AB, A'B')$ et la puissance P_0 dissipée par effet Joule dans cette ligne.
4. On ajoute en parallèle sur le dipôle D un condensateur de capacité $C = 130,4\,\mu\text{F}$.
- Calculer les intensités efficaces I_D dans le dipôle D , I_c dans le condensateur et I'_0 dans la ligne.
 - Déterminer la puissance P'_0 dissipée par effet Joule dans la ligne.
 - Comparer P_0 et P'_0 . Conclure quant à l'intérêt du condensateur.

Réponse fréquentielle d'un circuit linéaire - Filtres linéaires du premier et du second ordre

18

1. Notion de quadripôle

1.1 Quadripôle

Un *quadripôle* (Cf. figure 18.1) est un réseau électrique comportant quatre bornes de liaison avec l'extérieur :

- deux bornes reliées à un générateur d'entrée ou à un circuit linéaire représenté par son modèle de Thévenin ou de Norton ;
- deux bornes reliées à un circuit d'utilisation (aussi appelé circuit de charge).

Un quadripôle est dit linéaire s'il est composé uniquement de dipôles linéaires.

On choisira les conventions représentées sur la figure 18.2 :

- à l'entrée, le quadripôle est en convention récepteur et le générateur G en convention générateur ;
- à la sortie, le quadripôle est en convention générateur et la charge U en convention récepteur.

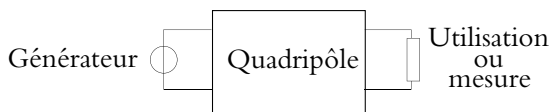


Figure 18.1 Schéma général d'un quadripôle.

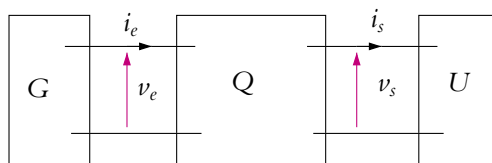


Figure 18.2 Conventions d'orientation.

► **Remarque :** Nous verrons dans la suite qu’il arrive très souvent que les deux bornes du bas soient reliées entre elles et reliées à la masse comme dans le cas où le quadripôle est un amplificateur opérationnel.

1.2 Complément : notion d’impédances d’entrée et de sortie

Si on considère un quadripôle en régime sinusoïdal, on peut adopter deux points de vue (Cf. figure 18.3). Le générateur ne voit qu’un dipôle passif, constitué en fait du quadripôle et de la charge ; ce dipôle est traversé par l’intensité i_e sous la tension v_e . La charge, elle, ne voit qu’un dipôle actif, constitué en fait du générateur et du quadripôle, générant une intensité i_s sous une tension v_s .

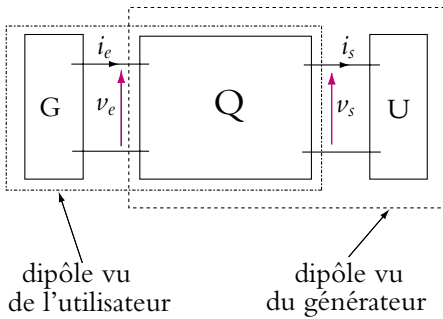


Figure 18.3 Quadripôle vu du générateur ou de l'utilisateur.

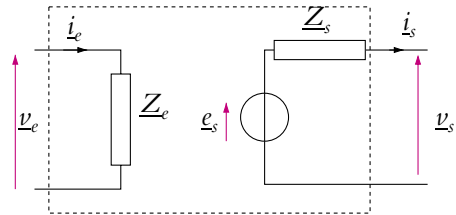


Figure 18.4 Modélisation des entrées et sorties d'un quadripôle.

On peut donc modéliser le quadripôle et sa charge en entrée par une impédance appelée *impédance d'entrée* et modéliser le générateur et le quadripôle en sortie par sa représentation de Thévenin ou Norton équivalente.

- L'impédance d'entrée $\underline{Z}_e$ est le rapport $\frac{v_e}{i_e}$. Elle est indépendante du générateur d'entrée mais peut dépendre de la charge (Cf. figure 18.4).
- Vu de la charge, le reste du circuit est donc un dipôle actif, on peut le représenter par son modèle de Thévenin ou de Norton dont l'impédance porte le nom d'impédance de sortie $\underline{Z}_s$ (Cf. figure 18.4).

2. Fonction de transfert

2.1 Transfert statique

On se place en régime continu et on désigne la grandeur d'entrée par E_c (intensité d'entrée I_e ou tension d'entrée U_e) et la grandeur de sortie par S_c (intensité de sortie I_s ou tension de sortie U_s). Le transfert statique est par définition :

$$H_0 = \frac{S_c}{E_c} \quad (18.1)$$

Dans le cas des rapports de tension ou d'intensité, on parle respectivement d'*amplification statique en tension et en courant*. Le rapport I_s/U_e est appelé transconductance statique et U_s/I_e est la transrésistance statique.

Par exemple, dans le cas de l'amplification statique en tension pour les quadripôles représentés sur les figures 18.5 et 18.6, les condensateurs se comportent comme des circuits ouverts puisque le courant est continu. Si l'on considère pour les deux cas qu'il n'y a pas de dipôle de charge, les courants dans les résistors sont nuls : pour le premier $U_s = 0$ V soit $H_0 = 0$ et pour le second $U_s = U_e$ soit $H_0 = 1$.

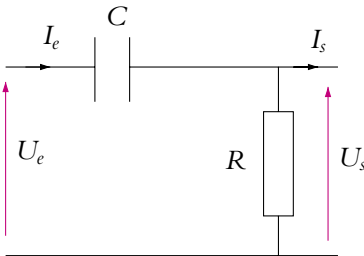


Figure 18.5 Transfert statique - Circuit CR.

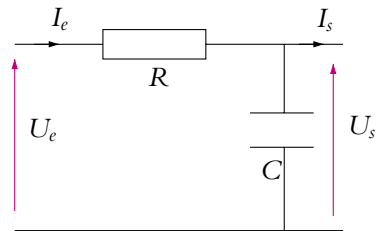


Figure 18.6 Transfert statique - Circuit RC.

2.2 Fonction de transfert en régime sinusoïdal forcé

On se place en régime sinusoïdal forcé et on désigne la grandeur d'entrée par e (intensité d'entrée i_e ou tension d'entrée u_e) et la grandeur de sortie par s (intensité de sortie i_s ou tension de sortie u_s). Le générateur d'entrée étant sinusoïdal :

$$e(t) = E\sqrt{2}\cos(\omega t)$$

en notant E la valeur efficace. La grandeur complexe associée s'écrit :

$$\underline{e}(t) = E\sqrt{2}\exp(j\omega t)$$

Dans un circuit linéaire, il a été vu que le générateur imposait sa pulsation donc en sortie le signal a même pulsation que le signal d'entrée :

$$s(t) = S\sqrt{2}\cos(\omega t + \varphi)$$

et en notation complexe :

$$\underline{s}(t) = S\sqrt{2}\exp(j(\omega t + \varphi)) = \underline{S}\sqrt{2}\exp(j\omega t) \quad \text{avec} \quad \underline{S} = S\exp(j\varphi)$$

On définit la *fonction de transfert complexe* par :

$$\underline{H}(j\omega) = \frac{\underline{s}}{\underline{e}} \quad (18.2)$$

C'est une fonction complexe dépendant de la pulsation. On peut la mettre sous la forme

$$\underline{H}(j\omega) = H(\omega) \exp(j\varphi(\omega)) \quad (18.3)$$

De l'équation précédente, il ressort que $H(\omega) = S/E$, rapport des valeurs efficaces, est le module de la fonction de transfert et que $\varphi(\omega)$, déphasage de la sortie par rapport à l'entrée, est son argument (défini à 2π près).

► **Remarque :**

1. $H(j\omega)$ est aussi le rapport des amplitudes des grandeurs $s(t)$ et $e(t)$.
2. Suivant les grandeurs de sortie et d'entrée, on parle d'amplification complexe en tension ou en courant, de transadmittance complexe et de transimpédance complexe. **On s'intéressera principalement aux cas où les grandeurs d'entrée et de sortie sont des tensions.**

Ainsi sur l'exemple de la figure 18.7, si on suppose que l'on branche en sortie un oscilloscope supposé idéal (d'impédance infinie), on peut appliquer un diviseur de tension pour calculer le rapport $\underline{u}_s/\underline{u}_e$:

$$\underline{H}(j\omega) = \frac{\underline{u}_s}{\underline{u}_e} = \frac{\underline{Z}_C}{R + \underline{Z}_C + \underline{Z}_L} = \frac{1}{1 + R\underline{Y}_C + \underline{Z}_L\underline{Y}_C} = \frac{1}{1 + jLC\omega - LC\omega^2}$$

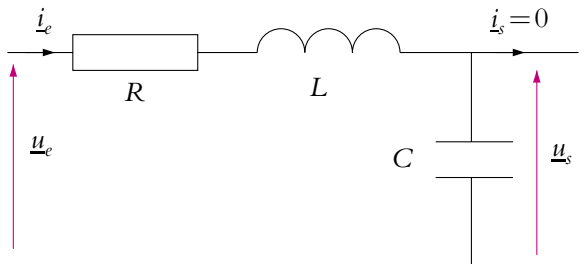


Figure 18.7 Fonction de transfert - Circuit RLC

2.3 Cas d'un quadripôle chargé

Dans le paragraphe précédent, la fonction de transfert a été calculée sur un quadripôle non chargé c'est-à-dire avec une charge d'impédance infinie et donc un courant de sortie nul. La fonction de transfert obtenue est appelée *fonction de transfert intrinsèque du quadripôle*.

Cependant, la charge n'est pas toujours d'impédance infinie et la fonction de transfert est alors modifiée.

On prend l'exemple suivant (le quadripôle est entouré par les pointillés) :

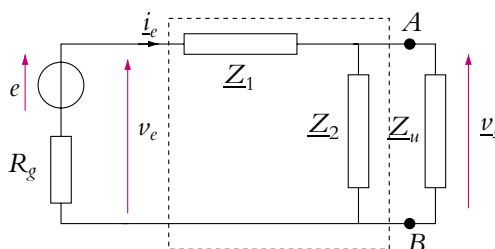


Figure 18.8 Influence de la charge sur la fonction de transfert.

On note $\underline{Z}_{eq}$, l'impédance équivalente à $\underline{Z}_2$ et $\underline{Z}_u$ en parallèle. On peut alors appliquer un diviseur de tension pour calculer la fonction de transfert :

$$\underline{H} = \frac{u_s}{u_e} = \frac{\underline{Z}_{eq}}{\underline{Z}_1 + \underline{Z}_{eq}}$$

La fonction de transfert obtenue est donc différente de la fonction de transfert intrinsèque ici égale à :

$$\frac{\underline{Z}_2}{\underline{Z}_1 + \underline{Z}_2}$$

$\underline{H}(j\omega)$ n'est pas une caractéristique intrinsèque du quadripôle : la fonction de transfert dépend de la charge.

En revanche, si l'impédance de la charge est très grande, $\underline{Z}_{eq} \simeq \underline{Z}_2$ et on retrouve la fonction de transfert intrinsèque.

3. Gain

Comme les variations de la fonction de transfert en fonction de la fréquence sont souvent importantes (facteur 10, 100 ou plus), on préfère utiliser une échelle logarithmique qu'on appelle le *gain en décibels*.

3.1 Gain en puissance

Si l'on utilise les notations de la figure 18.9, le quadripôle consomme à l'entrée la puissance moyenne $\langle P_e \rangle = U_e I_e \cos \varphi_e$, où U_e et I_e sont les valeurs efficaces d'entrée et φ_e l'argument de l'impédance d'entrée (égal au déphasage entre la tension et l'intensité).

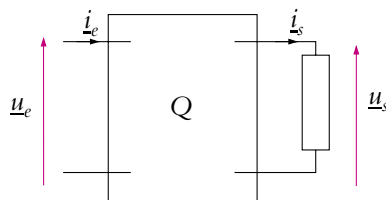


Figure 18.9 Gain en puissance.

De même, la charge reçoit la puissance moyenne $\langle P_s \rangle = U_s I_s \cos \varphi_s$ où U_s et I_s sont les valeurs efficaces de sortie et φ_s l'argument de l'impédance de la charge. Pour définir le gain en puissance, on utilise une échelle logarithmique :

$$G_p = \log \frac{\langle P_s \rangle}{\langle P_e \rangle} \quad (18.4)$$

On utilise le logarithme décimal plutôt que le logarithme népérien car il est plus facile de remonter aux rapport 10, 100, 1 000 ... entre la puissance de sortie et celle d'entrée. Le gain est une grandeur sans dimension ; cependant on a l'habitude de l'exprimer en bel, unité rendant en hommage à Graham Bell, l'inventeur du téléphone. En pratique, on utilise un sous-multiple du bel : le décibel (de symbole dB). On indique l'unité en indice de G :

$$G_{p(\text{dB})} = 10 \log \frac{\langle P_s \rangle}{\langle P_e \rangle}$$

3.2 Gain en tension

La puissance n'est pas la grandeur la plus facilement mesurable dans un circuit. Il a été dit précédemment qu'on s'intéresserait surtout aux tensions d'entrée et de sortie d'un circuit (grandeurs mesurables directement avec un oscilloscope). Comment étendre la définition du gain en puissance aux tensions ? Si on exprime les puissances moyennes en fonction de l'admittance d'entrée du quadripôle et en fonction de l'admittance de la charge :

$$\langle P_e \rangle = \text{Re}(\underline{Y}_e) U_e^2 \quad \text{et} \quad \langle P_s \rangle = \text{Re}(\underline{Y}_u) U_s^2$$

On peut alors écrire sous une autre forme le gain en puissance :

$$G_{p(\text{dB})} = 20 \log \frac{U_s}{U_e} + 10 \log \frac{\text{Re}(\underline{Y}_u)}{\text{Re}(\underline{Y}_e)}$$

Même si les parties réelles des admittances sont différentes, on garde comme définition du gain en tension en décibels :

$$G_{\text{dB}} = 20 \log \frac{U_s}{U_e} \quad (18.5)$$

Il faut savoir rapidement passer d'un gain en décibels à un gain en tension. Ainsi :

- $G_{\text{dB}} = 20 \text{ dB} \Leftrightarrow \frac{U_s}{U_e} = 10$
- $G_{\text{dB}} = 40 \text{ dB} \Leftrightarrow \frac{U_s}{U_e} = 100$
- $G_{\text{dB}} = -20 \text{ dB} \Leftrightarrow \frac{U_s}{U_e} = \frac{1}{10}$
- $G_{\text{dB}} = -40 \text{ dB} \Leftrightarrow \frac{U_s}{U_e} = \frac{1}{100}$

4. Étude du comportement en fonction de la fréquence. Diagramme de Bode

Pourquoi s'intéresse-t-on au comportement d'un quadripôle en fonction de la fréquence ? On a vu dans le chapitre sur l'étude du circuit R, L, C série que la tension aux bornes d'un des dipôles d'un circuit R, L, C série dépend de la fréquence.

Dans l'exemple d'une chaîne hi-fi, les enceintes comportent plusieurs haut-parleurs, l'un fonctionnant pour les basses fréquences, le deuxième pour les fréquences intermédiaires et le dernier pour les hautes fréquences. Il est nécessaire que l'amplificateur envoie les bonnes fréquences vers chacun des haut-parleurs. Pour cela, il possède des circuits appelés *filtres* laissant passer les basses fréquences et coupant les hautes ou bien laissant seulement passer les fréquences intermédiaires ou bien ne laissant passer que les hautes fréquences. On voit sur cet exemple qu'il est nécessaire de connaître le comportement fréquentiel des circuits.

4.1 Diagramme de Bode

Il faut choisir une représentation graphique qui permettra en un coup d'œil de se rendre compte du comportement du quadripôle quelle que soit la fréquence. Il en existe plusieurs dont la *représentation de Bode* que l'on va utiliser ici.

Il s'agit de l'une des représentations possibles de $\underline{H}(j\omega)$ en fonction de la pulsation ω . On choisit la pulsation plutôt que la fréquence car c'est elle qui apparaît naturellement dans les calculs en complexes comme cela a été vu précédemment. Comme $\underline{H}$ est un nombre complexe, on choisit de représenter le module sur un diagramme et l'argument (ou phase) sur un autre diagramme. **Un diagramme de Bode est donc en fait l'association de ces deux diagrammes.** Comme la pulsation est susceptible de varier sur un large intervalle de valeurs, on adopte une échelle logarithmique en abscisses. En effet, si on veut étudier le comportement du filtre entre 0 Hz et 1 MHz et que sur le diagramme on choisit une échelle de 10 cm pour 1 MHz, avec une échelle linéaire il faut prendre une très bonne loupe pour savoir ce qui se passe à 1 kHz correspondant à 0,1 mm alors qu'avec une échelle logarithmique, 1 kHz correspond à 5 cm puisque $\log 10^3 = 3$ et $\log 10^6 = 6$.

On représente le gain en décibels et la phase de $\underline{H}(j\omega)$ en fonction de $\log(\omega/\omega_{\text{ref}})$, ω_{ref} étant en général une pulsation caractéristique du circuit (ou de ω en échelle logarithmique si on dispose de papier semilogarithmique dont l'échelle en abscisses est logarithmique). Régulièrement on utilise la variable réduite sans dimension $x = \omega/\omega_{\text{ref}}$.

Un changement de pulsation de référence entraîne une translation horizontale du diagramme.

En travaux pratiques, on mesure directement des fréquences et non des pulsations. On représente alors le gain en fonction de la fréquence et non de la pulsation (en échelle logarithmique). Les deux diagrammes s'obtiennent l'un de l'autre par une translation horizontale de $\log_{10} 2\pi$.

Le fait de représenter le gain en décibels permet des constructions graphiques simples dans les cas où la fonction de transfert peut se mettre sous forme d'un produit de fonctions plus simples. Dans ce cas les gains en décibels s'ajoutent de même que les déphasages. En effet, si $\underline{H} = \underline{H}_1 \underline{H}_2$ avec :

$$\underline{H} = H \exp(j\varphi) \quad \text{et} \quad \underline{H}_1 = H_1 \exp(j\varphi_1), \quad \underline{H}_2 = H_2 \exp(j\varphi_2)$$

alors :

$$G = 20 \log H = 20 \log H_1 + 20 \log H_2 = G_1 + G_2 \quad \text{et} \quad \varphi = \varphi_1 + \varphi_2$$

Il est alors relativement aisé d'effectuer graphiquement une somme de deux courbes. Ce point sera détaillé au paragraphe 5.4.

4.2 Bande passante

a) Définition

Lorsque le gain G_{dB} présente un maximum noté $G_{\text{dB,max}}$, on définit la (les) pulsation(s) de coupure à -3 dB , ω_c telles que :

$$G_{\text{dB}}(\omega_c) = G_{\text{dB,max}} - 3 \text{ dB} \quad (18.6)$$

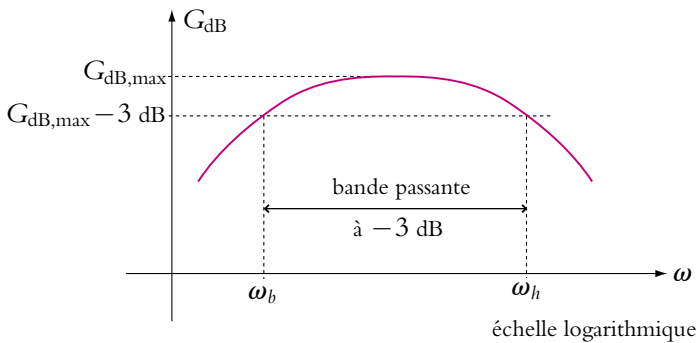


Figure 18.10 Bande passante.

On définit de même la (ou les) fréquence(s) de coupure par $f_c = \frac{\omega_c}{2\pi}$. Dans le cas représenté sur la figure 18.10, il y a deux pulsations de coupure :

- ω_b pulsation de coupure basse,

- ω_h pulsation de coupure haute.

On appelle bande passante, l'intervalle de pulsation $[\omega_b, \omega_h]$ pour lequel

$$G_{dB} \geq G_{dB,max} - 3 \text{ dB}$$

On note $\Delta\omega$ la largeur de la bande passante :

$$\Delta\omega = \omega_h - \omega_b \quad (18.7)$$

Par abus de langage, on utilise fréquemment le terme « bande passante » pour désigner $\Delta\omega$.

► **Remarque :** On peut définir aussi la bande passante avec les fréquences de coupure à la place des pulsations. C'est ce qu'on fait en général en travaux pratiques puisqu'on mesure des fréquences et non des pulsations. Dans ce cas, la largeur de la bande passante est $\Delta f = f_h - f_b$ où f_b et f_h sont respectivement les fréquences de coupure basse et haute.

Il arrive qu'il n'y ait qu'une pulsation de coupure ω_c comme dans les cas des figures 18.11 et 18.12.

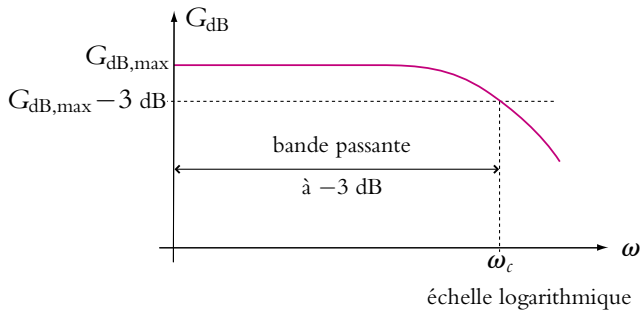


Figure 18.11 Pulsation de coupure haute et bande passante.

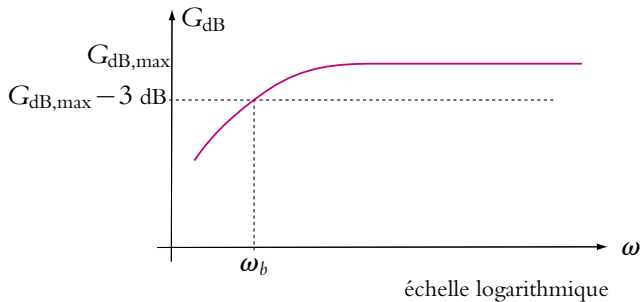


Figure 18.12 Pulsation de coupure basse.

Dans le premier cas, le gain est supérieur à $G_{\text{dB,max}} - 3 \text{ dB}$ pour les pulsations inférieures à ω_c . La pulsation de coupure ω_c est donc une pulsation de coupure haute. Puisqu'il n'y a pas de pulsation de coupure basse, la bande passante est $[0, \omega_c]$ et la largeur de bande passante :

$$\Delta\omega = \omega_c - 0 = \omega_c$$

Dans le second cas, le gain est supérieur à $G_{\text{dB,max}} - 3 \text{ dB}$ pour les pulsations supérieures à ω_c . La pulsation de coupure ω_c est donc une pulsation de coupure basse. La bande passante est $[\omega_c, +\infty[$ et la largeur de bande passante n'est pas définie.

b) Signification physique

Que représente physiquement la bande passante ?

On revient à l'exemple de la figure 18.10. Si un signal d'entrée sinusoïdal envoyé sur le quadripôle a sa pulsation ω dans la bande passante du quadripôle ($\omega_b \leq \omega \leq \omega_h$), son amplification aura une valeur proche de celle de l'amplification maximale. Au contraire, pour un signal dont la pulsation est hors bande passante, l'amplification sera plus faible. Ainsi, en sortie, un signal hors bande passante sera assez fortement atténué par rapport à un signal dans la bande passante et de même amplitude en entrée, d'où le nom de bande passante.

On considérera qu'un filtre coupe les composantes du signal dont la pulsation est hors bande passante et laisse passer celles dont la pulsation est dans la bande passante.

À quelle valeur de la fonction de transfert correspond la valeur « -3 dB » ? Si l'on cherche x tel que $20 \log x = 3$, on trouve $x = \sqrt{2}$, soit

$$G_{\text{dB}}(\omega) = G_{\text{dB,max}} - 3 \Leftrightarrow 20 \log H = 20 \log \frac{H_{\text{max}}}{\sqrt{2}} \Leftrightarrow H = \frac{H_{\text{max}}}{\sqrt{2}}$$

Il s'agit donc de la même définition que celle adoptée pour la bande passante de la résonance d'intensité du circuit R, L, C série.

5. Filtres du premier ordre

On appelle *circuits ou filtres du premier ordre* les circuits dont l'équation différentielle entre la tension d'entrée et la tension de sortie est du premier ordre. La fonction de transfert ne contient que des termes de puissance au plus un en ω .

Dans la suite de ce chapitre, on s'intéresse uniquement à la fonction de transfert intrinsèque de quadripôles (avec une charge d'impédance infinie).

5.1 Filtre passe-bas du premier ordre

L'étude sera complètement détaillée pour ce premier filtre et plus rapide pour les suivants. On étudie le quadripôle de la figure 18.13.

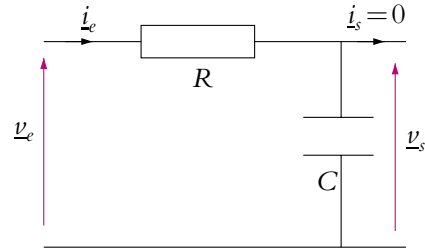


Figure 18.13 Exemple de filtre passe-bas du premier ordre.

a) Étude rapide

Une étude de filtre doit toujours commencer par l'étude du comportement en très basses fréquences (T.B.F.) lorsque $\omega \rightarrow 0$ et en très hautes fréquences (T.H.F.) lorsque $\omega \rightarrow \infty$ pour connaître rapidement la nature du filtre.

En T.B.F., les condensateurs se comportent comme des interrupteurs ouverts et en T.H.F. comme des fils d'où les circuits équivalents au filtre en T.B.F. et en T.H.F. représentés sur les figures 18.14 et 18.15.

En T.B.F., le courant est nul dans R et donc $v_s = v_e$ et $H = 1$ (Cf. figure 18.14).

En T.H.F., le condensateur se comporte comme un court-circuit donc $v_s = 0$ V et $H = 0$ (Cf. figure 18.15).

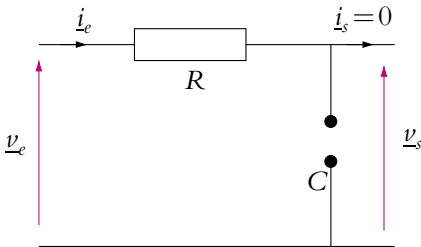


Figure 18.14 Circuit équivalent en basses fréquences.

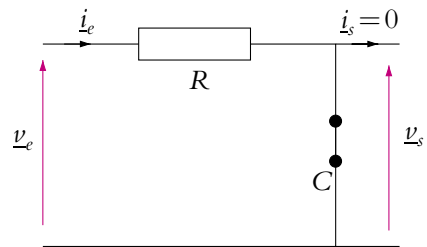


Figure 18.15 Circuit équivalent en hautes fréquences.

Il s'agit donc d'un circuit qui laisse passer les signaux de basses fréquences et atténue ceux de hautes fréquences, on dit que ce filtre est un *filtre passe-bas*.

b) Calcul de la fonction de transfert

Étant donné que le courant de sortie est nul, c'est le même courant qui circule dans les deux dipôles qui sont donc en série. On peut appliquer un diviseur de tension :

$$H = \frac{v_s}{v_e} = \frac{Z_c}{R + Z_c} = \frac{1}{1 + RY_c} = \frac{1}{1 + jRC\omega}$$

On constate que ce circuit ne fait intervenir qu'un terme de puissance au plus un en ω . Le module de la fonction de transfert vaut :

$$H = \frac{1}{\sqrt{1 + (RC\omega)^2}} \quad (18.8)$$

On vérifie que l'étude rapide donne des résultats cohérents avec la fonction trouvée ci-dessus :

- $H \rightarrow 0$ lorsque $\omega \rightarrow \infty$,
- $H \rightarrow 1$ lorsque $\omega \rightarrow 0$.

On voit immédiatement que H est une fonction décroissante de ω (il est inutile de calculer la dérivée). Le maximum de H est donc atteint lorsque ω tend vers 0 et vaut 1 : $H_{\max} = 1$.

c) Pulsation de coupure

On cherche les pulsations ω_c telles que $H(\omega_c) = \frac{H_{\max}}{\sqrt{2}} = \frac{1}{\sqrt{2}}$.

Elles vérifient $1 + (RC\omega_c)^2 = 2$ soit $\omega_c = \pm(RC)^{-1}$. Or physiquement, seules les pulsations positives ont un sens. Ce filtre n'a qu'une pulsation de coupure haute :

$$\omega_c = \frac{1}{RC} \quad (18.9)$$

Il n'y a pas de pulsation de coupure basse et la plus petite pulsation est 0 rad.s^{-1} . La largeur de la bande passante est (Cf. paragraphe 4.2) :

$$\Delta\omega = \omega_c - 0 = \omega_c = \frac{1}{RC} \quad (18.10)$$

d) Équation canonique d'un filtre passe-bas du premier ordre

La fonction de transfert statique n'est pas toujours égale à 1. De manière générale, elle vaut H_0 . Avec l'expression de la pulsation de coupure déterminée au paragraphe précédent, on peut donc écrire l'expression générale de la fonction de transfert d'un filtre passe-bas du premier ordre :

$$\underline{H}(j\omega) = \frac{H_0}{1 + j\frac{\omega}{\omega_c}} \quad (18.11)$$

Il semble naturel de choisir ici $\omega_{\text{ref}} = \omega_c$ et la variable réduite est $x = \frac{\omega}{\omega_c}$. Dans ce cas, la fonction de transfert s'écrit :

$$\underline{H}(jx) = \frac{H_0}{1 + jx} \quad (18.12)$$

Cette expression, qu'il faut connaître, est appelée *expression canonique des filtres passe-bas du premier ordre*.

Dans le cas où le filtre est réalisé avec d'autres composants (R et L) par exemple, la forme canonique reste la même mais les expressions de H_0 et ω_c changent.

e) Étude du gain en dB

De l'équation canonique (18.12), on tire le gain en décibels :

$$G_{dB} = 20 \log \frac{|H_0|}{\sqrt{1+x^2}} = 20 \log |H_0| - 10 \log(1+x^2)$$

Dans la suite, on note G_0 le terme $20 \log |H_0|$.

Quel est le comportement asymptotique du gain en T.B.F. et T.H.F. ?

En T.B.F., on peut écrire $\omega \ll \omega_c$ soit $1 \gg x (= \frac{\omega}{\omega_c})$. On peut donc négliger le terme en x dans G_{dB} . Puisque $\log 1 = 0$, on trouve l'équation d'une première asymptote (notée G_{dB}^{a1}) pour la fonction G_{dB} :

$$G_{dB}^{a1} = G_0 \quad \text{en T.B.F.}$$

En T.H.F., c'est le contraire : $\omega \gg \omega_c$ soit $1 \ll x (= \frac{\omega}{\omega_c})$. On garde donc uniquement le terme en x dans G_{dB} . On trouve alors l'équation d'une deuxième asymptote (notée G_{dB}^{a2}) :

$$G_{dB}^{a2} = G_0 - 20 \log x = G_0 - 20 \log \frac{\omega}{\omega_c} \quad \text{en T.H.F.}$$

On veut tracer G_{dB} en fonction de $\log x$. L'équation de G_{dB}^{a2} est l'équation d'une droite en coordonnées logarithmiques. En effet, l'ordonnée Y est égale à $G_{dB}^{a2}(\omega)$, et l'abscisse X à $\log x = \log(\omega/\omega_c)$. On obtient bien l'équation d'une droite :

$$Y = G_0 - 20X$$

Quelle est la pente de cette droite ?

On calcule la valeur de l'asymptote pour deux pulsations ω_1 et $\omega_2 = 10 \omega_1$, donc pour deux valeurs de x , x_1 et $x_2 = 10 x_1$:

$$G_{dB}^{a2}(x_1) = G_0 - 20 \log x_1$$

$$G_{dB}^{a2}(x_2) = G_0 - 20 \log x_2 = G_0 - 20 \log 10 - 20 \log x_1$$

Soit $G_{dB}^{a2}(x_2) - G_{dB}^{a2}(x_1) = -20$ dB. Ainsi lorsque la pulsation est multipliée par 10 (on parle d'une décade), le gain est diminué de 20 dB, on dit que l'asymptote a une pente de **-20 dB/décade** (souvent noté -20 dB/déc).

► **Remarque :** Lorsque la pulsation est multipliée par deux (les musiciens utilisent alors le terme d'octave), le gain est diminué de 6 dB. La pente de l'asymptote est donc aussi de -6 dB/octave.

En quel point se croisent les asymptotes ?

On écrit l'égalité des équations des deux droites :

$$G_{\text{dB}}^{a_2}(x) = G_{\text{dB}}^{a_1}(x) \Leftrightarrow 20 \log x = 0 \Leftrightarrow x = 1 \Leftrightarrow \omega = \omega_c$$

Avant de tracer la courbe, on détermine sa position par rapport aux asymptotes :

$$G_{\text{dB}}(x) - G_{\text{dB}}^{a_1}(x) = 20 \log \frac{1}{\sqrt{1+x^2}} < 0$$

$$G_{\text{dB}}(x) - G_{\text{dB}}^{a_2}(x) = 20 \log \frac{x}{\sqrt{1+x^2}} < 0$$

Dans les deux cas, la courbe est au-dessous de l'asymptote.

On peut déterminer numériquement l'écart entre la courbe et les asymptotes :

- Pour $x = 0,1$, soit $\omega = 0,1 \omega_c$,
 $G_{\text{dB}} = -4,3 \cdot 10^{-3}$ dB alors que l'asymptote vaut 0 dB,
- Pour $x = 10$, soit $\omega = 10 \omega_c$,
 $G_{\text{dB}} = -20,04$ dB alors que l'asymptote vaut -20 dB.

Ainsi, sauf pour des valeurs de ω proche de ω_c , la courbe est quasiment confondue avec les asymptotes.

Le diagramme est représenté sur la figure 18.16 pour trois valeurs de H_0 : 0,1 ; 1 (cas du circuit de la figure 18.13) et 10. On commence toujours le tracé par celui des asymptotes avant celui de la courbe réelle. On remarque que c'est à la pulsation de coupure que l'écart entre la courbe et les asymptotes est le plus grand.

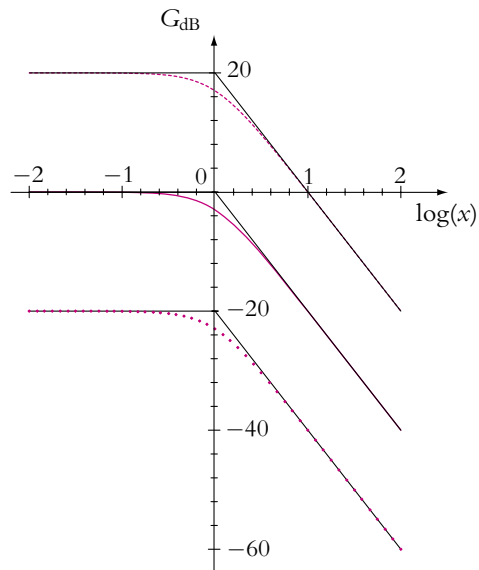


Figure 18.16 Diagramme de Bode en gain d'un filtre passe-bas du premier ordre pour différentes valeurs de H_0 : points : $H_0 = 0,1$, trait plein : $H_0 = 1$, pointillés : $H_0 = 10$.

f) Étude de la phase en fonction de la pulsation

Il est intéressant de calculer les équivalents de la fonction de transfert pour déterminer rapidement les phases en T.B.F. et T.H.F. Ainsi en T.B.F., $\underline{H}$ est équivalent à H_0 ; c'est un réel d'argument $\text{Arg } H_0$.

En T.H.F, $\underline{H}$ est équivalent à $\frac{H_0}{jx} = -jH_0 \frac{\omega_c}{\omega}$. C'est un imaginaire pur d'argument $-\frac{\pi}{2} + \text{Arg } H_0$.

Dans le cas général où H_0 est différent de 1, il faut faire attention à son signe :

- si H_0 est positif, $\text{Arg } H_0 = 0$;
- si H_0 est négatif, $\text{Arg } H_0 = \pi$.

On fait l'hypothèse dans la suite que H_0 est positif. Dans le cas contraire, il faut penser à ajouter π aux résultats obtenus.

Pour faire une étude plus précise, on utilise le fait que l'argument φ est égal à la différence entre les arguments du numérateur (φ_N) et du dénominateur (φ_D). Soit ici $\varphi = -\varphi_D$. Or :

$$\tan \varphi_D = \frac{\omega}{\omega_c}$$

Est-ce suffisant pour écrire que $\varphi_D = \arctan \frac{\omega}{\omega_c}$? Non, il y a une indétermination sur φ_D car φ_D et $\varphi_D + \pi$ ont même tangente. Il faut déterminer l'intervalle de longueur π dans lequel se trouve φ_D . Pour cela, on regarde, par exemple, le signe de la partie réelle de H qui est aussi le signe de $\cos \varphi_D$. On peut aussi déterminer le signe de la partie imaginaire si c'est plus simple ; il est de même signe que $\sin \varphi_D$. Ici $\cos \varphi_D$ est positif donc $-\frac{\pi}{2} \leq \varphi_D \leq \frac{\pi}{2}$. On peut alors écrire :

$$\varphi = -\varphi_D = -\arctan x = -\arctan \frac{\omega}{\omega_c}$$

On retrouve les comportements T.B.F. et T.H.F. déjà déterminés.

On note que pour $x = 1$ (soit $\omega = \omega_c$), $\varphi = -\frac{\pi}{4}$.

On obtient le graphe de la figure 18.17. On peut aussi tracer deux diagrammes asymptotiques représentés sur la figure : l'un en marche d'escalier et l'autre reliant les deux asymptotes horizontales par une droite oblique entre $\omega_c/10$ et $10 \omega_c$.

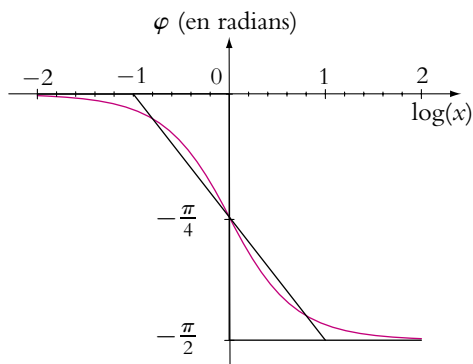


Figure 18.17 Diagramme de Bode en phase d'un filtre passe-bas du premier ordre.

g) Comportement intégrateur

Comme cela a été vu au paragraphe précédent, pour $\omega \gg \omega_c$, c'est-à-dire lorsque l'on peut confondre la courbe du gain avec l'asymptote, la fonction de transfert peut s'écrire :

$$\underline{H} = \frac{H_0 \omega_c}{j\omega}$$

or : $\underline{H} = \underline{v_s} / \underline{v_e}$. En combinant ces deux expressions, on obtient :

$$\underline{v_s} = H_0 \omega_c \frac{\underline{v_e}}{j\omega}$$

On se rappelle qu'avec des signaux complexes, diviser par $j\omega$ revient à intégrer. On en déduit qu'un filtre possédant une asymptote de -20 dB/déc se comporte dans la zone de l'asymptote comme un montage intégrateur avec :

$$\underline{v_s} = H_0 \omega_c \frac{\underline{v_e}}{j\omega} \quad \text{soit} \quad v_s = H_0 \omega_c \int v_e dt$$

5.2 Filtre passe-haut du premier ordre

On étudie maintenant le filtre de la figure 18.18.

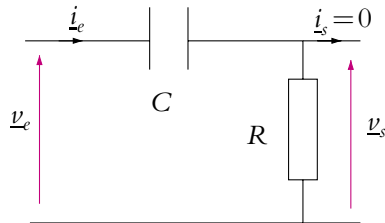


Figure 18.18 Exemple de filtre passe-haut du premier ordre.

a) Étude rapide

En T.B.F., les condensateurs se comportent comme des interrupteurs ouverts, et en T.H.F. comme des fils, d'où les deux circuits équivalents représentés sur les figures 18.19 et 18.20.

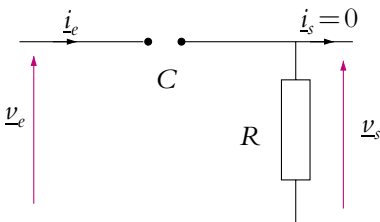


Figure 18.19 Circuit équivalent en basses fréquences.

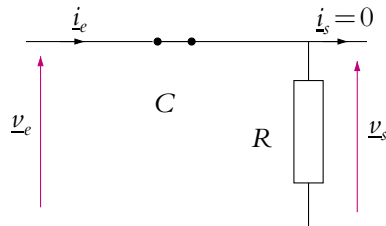


Figure 18.20 Circuit équivalent en hautes fréquences.

En T.B.F., le courant est nul dans R et donc $\underline{v}_s = 0$ V et $\underline{H} = 0$ (Cf. figure 18.19).

En T.H.F., le condensateur se comporte comme un court-circuit donc $\underline{v}_s = \underline{v}_c$ et $\underline{H} = 1$ (Cf. figure 18.20).

Il s'agit donc d'un circuit qui laisse passer les signaux de hautes fréquences et atténue ceux de basses fréquences, on dit que ce filtre est *un filtre passe-haut*.

b) Calcul de la fonction de transfert

Étant donné que le courant de sortie est nul, c'est le même courant qui circule dans les deux dipôles qui sont donc en série. On peut appliquer un diviseur de tension :

$$\underline{H} = \frac{\underline{v}_s}{\underline{v}_c} = \frac{R}{R + \underline{Z}_c} = \frac{\underline{Y}_c}{1 + R\underline{Y}_c} = \frac{jRC\omega}{1 + jRC\omega}$$

On peut écrire aussi $\underline{H}$ sous la forme :

$$\underline{H} = \frac{1}{1 + \frac{1}{jRC\omega}}$$

On constate que ce circuit ne fait intervenir qu'un terme de puissance au plus un en ω . Le module de la fonction de transfert est :

$$H = \frac{1}{\sqrt{1 + \frac{1}{(RC\omega)^2}}} \quad (18.13)$$

On vérifie que l'étude rapide donne des résultats cohérents avec la fonction trouvée ci-dessus :

- $H \rightarrow 1$ lorsque $\omega \rightarrow \infty$,
- $H \rightarrow 0$ lorsque $\omega \rightarrow 0$.

On voit que H est une fonction croissante de ω . Le maximum de H est donc atteint lorsque $\omega \rightarrow \infty$ et vaut 1 : $H_{\max} = 1$.

c) Pulsation de coupure

On cherche les pulsations ω_c telles que $H(\omega_c) = \frac{H_{\max}}{\sqrt{2}} = \frac{1}{\sqrt{2}}$. Elles vérifient $1 + (RC\omega_c)^{-2} = 2$ soit $\omega_c = \pm(RC)^{-1}$. Comme précédemment, on ne garde que la solution positive. Ce filtre n'a donc qu'une pulsation de coupure basse :

$$\omega_c = \frac{1}{RC} \quad (18.14)$$

d) Équation canonique d'un filtre passe-haut du premier ordre

En haute fréquence, la fonction de transfert n'est pas toujours égale à 1. De manière générale elle vaut H_0 . Avec l'expression de la pulsation de coupure déterminée au paragraphe précédent on peut donc écrire l'expression générale de la fonction de transfert d'un filtre passe-haut du premier ordre :

$$\underline{H}(j\omega) = \frac{H_0}{1 - j\frac{\omega_c}{\omega}} = \frac{H_0 \frac{j\omega}{\omega_c}}{1 + j\frac{\omega}{\omega_c}} \quad (18.15)$$

Cette expression, qu'il faut connaître, est la *forme canonique des filtres passe-haut du premier ordre*.

Comme pour le filtre passe-bas, on choisit naturellement $\omega_{\text{ref}} = \omega_c$ et pour variable réduite $x = \frac{\omega}{\omega_c}$. L'équation canonique avec la variable x est donc :

$$\underline{H}(j\omega) = \frac{H_0}{1 - j\frac{1}{x}} = \frac{H_0 jx}{1 + jx} \quad (18.16)$$

Comme pour le filtre passe-bas, si le filtre est réalisé avec d'autres composants (R et L) par exemple, la forme canonique reste la même mais les expressions de H_0 et ω_c peuvent changer.

e) Étude du gain en dB

D'après l'expression canonique, le gain en décibels est :

$$G_{\text{dB}} = 20 \log \frac{H_0}{\sqrt{1 + \frac{1}{x^2}}} = G_0 - 10 \log \left(1 + \frac{1}{x^2} \right)$$

avec $G_0 = 20 \log H_0$.

Quel est le comportement asymptotique du gain en T.B.F. et T.H.F. ?

En T.B.F., $\omega \ll \omega_c$ soit $1 \gg x \left(= \frac{\omega}{\omega_c} \right)$. On garde donc uniquement le terme en x dans G_{dB} . On trouve alors l'équation d'une première asymptote :

$$G_{\text{dB}}^{a_1} = G_0 - 20 \log \frac{1}{x} = G_0 + 20 \log x = G_0 + 20 \log \frac{\omega}{\omega_c} \quad \text{en T.B.F.}$$

En s'inspirant de l'étude du filtre passe-bas, on peut dire que l'expression précédente est l'équation d'une droite mais cette fois-ci de pente $+20$ dB/déc.

En T.H.F., on peut écrire $\omega \gg \omega_c$ soit $1 \ll x \left(= \frac{\omega}{\omega_c} \right)$. On peut donc négliger le terme en x dans G_{dB} . Puisque $\log 1 = 0$, on trouve l'équation d'une deuxième asymptote :

$$G_{\text{dB}}^{a_2} = G_0 \quad \text{en T.H.F.}$$

En quel point se croisent les asymptotes ?

On écrit l'égalité des équations des deux droites :

$$G_{dB}^{a_2}(x) = G_{dB}^{a_1}(x) \Leftrightarrow 20 \log x = 0$$

$$\Leftrightarrow x = 1 \Leftrightarrow \omega = \omega_c$$

On laisse au lecteur le soin de montrer (pour s'exercer) que la courbe est au-dessous de ses asymptotes.

Comme pour le filtre passe-bas, la courbe est quasiment confondue avec les asymptotes sauf pour des valeurs de ω proches de ω_c .

Le diagramme est représenté sur la figure 18.21 pour trois valeurs de H_0 : 0,1 ; 1 (cas du circuit de la figure 18.18) et 10. On rappelle que le tracé commence toujours par celui des asymptotes avant celui de la courbe réelle.

On remarque que c'est encore à la pulsation de coupure que l'écart entre la courbe et les asymptotes est le plus grand.

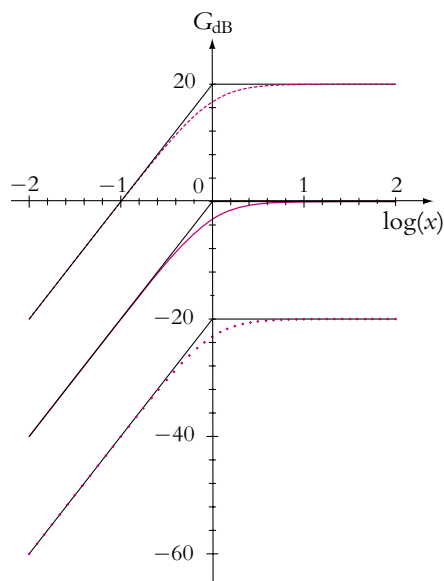


Figure 18.21 Diagramme de Bode en gain d'un filtre passe-haut du premier ordre pour différentes valeurs de H_0 :
points : $H_0 = 0,1$, trait plein :
 $H_0 = 1$, pointillés : $H_0 = 10$

f) Étude de la phase en fonction de la pulsation

Comme pour le filtre passe-bas, il est intéressant de calculer les équivalents de la fonction de transfert pour déterminer rapidement les phases en T.B.F. et T.H.F.

En T.B.F., $\underline{H}$ est équivalent à $\frac{H_0}{-j\frac{1}{x}} = j\frac{H_0 \omega}{\omega_c}$. C'est un imaginaire pur d'argument $\frac{\pi}{2} + \text{Arg } H_0$.

En T.H.F., $\underline{H}$ est équivalent à H_0 ; c'est un réel d'argument $\text{Arg } H_0$.

Comme pour le filtre passe-bas, dans le cas général où H_0 est différent de 1, il faut faire attention à son signe :

- si H_0 est positif, $\text{Arg } H_0 = 0$;
- si H_0 est négatif, $\text{Arg } H_0 = \pi$.

On fait l'hypothèse que H_0 est positif. Dans le cas contraire, il faut penser à ajouter π aux résultats obtenus.

Pour faire une étude plus précise, on utilise toujours le fait que $\varphi = \varphi_N - \varphi_D$. Or :

$$\tan \varphi_D = -\frac{1}{x} = -\frac{\omega_c}{\omega}$$

Est-ce suffisant pour écrire que $\varphi_D = \arctan \frac{1}{x}$? Comme précédemment, il y a une indétermination sur φ_D . On détermine l'intervalle de longueur π dans lequel se trouve φ_D . Ici $\cos \varphi_D$ est positif donc $-\frac{\pi}{2} \leq \varphi_D \leq \frac{\pi}{2}$. On peut alors écrire :

$$\varphi = -\varphi_D = \arctan x = +\arctan \frac{\omega_c}{\omega_\varphi}$$

On retrouve les comportements T.B.F. et T.H.F. déjà déterminés.

On note que pour $\omega = \omega_c$, $\varphi = \frac{\pi}{4}$.

On obtient le graphe de la figure 18.22.

On peut aussi tracer les deux diagrammes asymptotiques : celui en marche d'escalier et l'autre reliant les deux asymptotes horizontales par une droite oblique entre $\omega_c/10$ et $10\omega_c$.

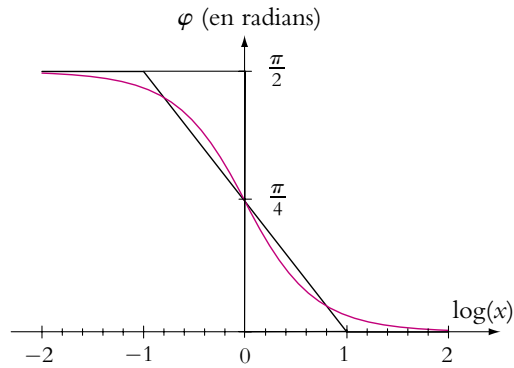


Figure 18.22 Diagramme de Bode en phase d'un filtre passe-haut du premier ordre.

g) Comportement dérivateur

Comme cela a été vu au paragraphe précédent, pour $\omega \ll \omega_c$, c'est-à-dire lorsque l'on peut confondre la courbe du gain avec l'asymptote, la fonction de transfert peut s'écrire :

$$\underline{H} = \frac{jH_0 \omega}{\omega_c}$$

Or : $\underline{H} = \frac{\underline{v}_s}{\underline{v}_e}$, en combinant ces deux expressions, on obtient :

$$\underline{v}_s = jH_0 \omega \frac{\underline{v}_e}{\omega_c}$$

On se rappelle qu'avec des signaux complexes, multiplier par $j\omega$ revient à dériver. On en déduit donc qu'un filtre possédant une asymptote de +20 dB/déc se comporte dans la zone de l'asymptote comme un montage dérivateur avec :

$$\underline{v}_s = \frac{H_0}{\omega_c} j\omega \underline{v}_e \quad \text{soit} \quad v_s = \frac{H_0}{\omega_c} \frac{dv_e}{dt}$$

5.3 Récapitulatif

Avec un filtre d'ordre 1, on ne peut avoir que des asymptotes de pentes 0 dB/déc et -20 dB/déc ou 20 dB/déc.

En ce qui concerne la phase, elle ne peut varier que de $\pi/2$ en valeur absolue et sa variation est monotone (soit croissante, soit décroissante).

Si l'on désire des variations de phase plus importantes ou des pentes plus importantes, il est nécessaire de combiner plusieurs filtres d'ordre 1 ou d'utiliser des filtres d'ordre supérieur qui seront étudiés dans la suite.

5.4 Composition de filtres d'ordre 1

Il arrive qu'une fonction de transfert puisse être mise sous la forme de produits ou rapports de fonction de transferts d'ordre 1. On va montrer sur un exemple qu'il est alors aisé de tracer les diagrammes de Bode en gain et en phase en sommant les courbes graphiquement, comme on l'a déjà signalé au paragraphe 4.1.

Soit $\underline{H} = \frac{1 + j\frac{\omega}{\omega_1}}{1 + j\frac{\omega}{\omega_2}}$, avec $\omega_1 > \omega_2$. On peut alors décomposer $\underline{H}$ en un produit de deux fonctions $\underline{H}_1$ et $\underline{H}_2$ avec :

$$\underline{H}_1 = 1 + j\frac{\omega}{\omega_1} \quad \text{et} \quad \underline{H}_2 = \frac{1}{1 + j\frac{\omega}{\omega_2}}$$

On peut alors étudier séparément les deux fonctions de transfert.

a) Diagramme en gain

On écrit le gain en décibels comme une somme :

$$G_{\text{dB}} = 20 \log H = 20 \log H_1 + 20 \log H_2 = G_{1,\text{dB}} + G_{2,\text{dB}}$$

► Étude asymptotique de $G_{1,\text{dB}} = 20 \log \sqrt{1 + \left(\frac{\omega}{\omega_1}\right)^2}$:

- Pour $\omega \ll \omega_1$, $\frac{\omega}{\omega_1} \ll 1$. L'asymptote est $G_{1,\text{dB}}^{\text{a1}} = 0$ dB.
- Pour $\omega \gg \omega_1$, $\frac{\omega}{\omega_1} \gg 1$. L'asymptote est $G_{1,\text{dB}}^{\text{a2}} = 20 \log \frac{\omega}{\omega_1}$, soit une droite de pente 20 dB/déc.

Les deux droites se croisent pour $\omega = \omega_1$.

► Étude asymptotique de $G_{2,\text{dB}} = -20 \log \sqrt{1 + \left(\frac{\omega_2}{\omega}\right)^2}$:

- Pour $\omega \ll \omega_2$, $\frac{\omega_2}{\omega} \gg 1$. L'asymptote est $G_{2,\text{dB}}^{a_1} = 20 \log \frac{\omega}{\omega_2}$, soit une droite de pente +20 dB/déc.
- Pour $\omega \gg \omega_2$, $\frac{\omega_2}{\omega} \ll 1$. L'asymptote est $G_{2,\text{dB}}^{a_2} = 0$ dB.

Les deux droites se croisent pour $\omega = \omega_2$.

On effectue le tracé des quatre asymptotes puis pour avoir les asymptotes de la courbe finale on ajoute les droites (Cf. figure 18.23). Il y a trois zones :

- Zone 1 pour $\omega < \omega_2$: somme de $G_{2,\text{dB}}^{a_1}$ et de $G_{1,\text{dB}}^{a_1}$, soit une droite de pente +20 dB/déc.
- Zone 2 pour $\omega_2 < \omega < \omega_1$: somme de $G_{2,\text{dB}}^{a_2}$ et de $G_{1,\text{dB}}^{a_1}$, soit une droite de pente +0 dB/déc.
- Zone 3 pour $\omega_1 < \omega$: somme de $G_{2,\text{dB}}^{a_2}$ et de $G_{1,\text{dB}}^{a_2}$, soit une droite de pente +20 dB/déc.

Il suffit alors de tracer la courbe réelle qui ici traverse l'asymptote.

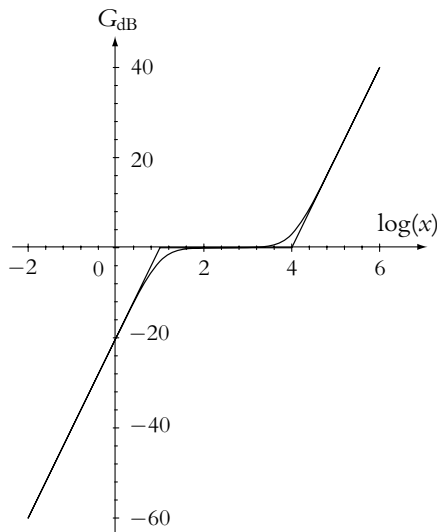


Figure 18.23 Composition de diagrammes de Bode en gain.

avec $x = \frac{\omega}{\omega_{\text{ref}}}$ où $\omega_{\text{ref}} = 1 \text{ rad} \cdot \text{s}^{-1}$.

b) Diagramme en phase

La phase φ est la somme des phases φ_1 et φ_2 . On commence donc par étudier chaque phase séparément. Sachant que la phase pour chaque fonction d'ordre 1 varie de

$\pi/2$ en valeur absolue et qu'elle varie de manière monotone, il suffit de trouver les comportements T.B.F. et T.H.F.

► Étude de la phase de $\underline{H}_1 = 1 + j\frac{\omega}{\omega_1}$:

- Pour $\omega \ll \omega_1$, $\frac{\omega}{\omega_1} \ll 1$, $\underline{H}_1 \simeq 1$ et l'argument est $\varphi_1 = 0$.
- Pour $\omega \gg \omega_1$, $\frac{\omega}{\omega_1} \gg 1$, $\underline{H}_1 \simeq j\frac{\omega}{\omega_1}$ et l'argument est $\varphi_1 = \pi/2$.

D'autre part, $\varphi_1 = \pi/4$ pour $\omega = \omega_1$.

► Étude de la phase de $\underline{H}_2 = \left(1 + j\frac{\omega_2}{\omega}\right)^{-1}$:

- Pour $\omega \ll \omega_2$, $\frac{\omega}{\omega_2} \ll 1$, $\underline{H}_2 \simeq \left(j\frac{\omega_2}{\omega}\right)^{-1}$ et l'argument est $\varphi_2 = -\pi/2$.
- Pour $\omega \gg \omega_2$, $\frac{\omega}{\omega_2} \gg 1$, $\underline{H}_2 \simeq 1$ et l'argument est $\varphi_2 = 0$.

D'autre part, $\varphi_2 = -\frac{\pi}{4}$ pour $\omega = \omega_2$.

On trace alors φ_1 et φ_2 (Cf. figure 18.24) puis on effectue graphiquement la somme pour obtenir φ .

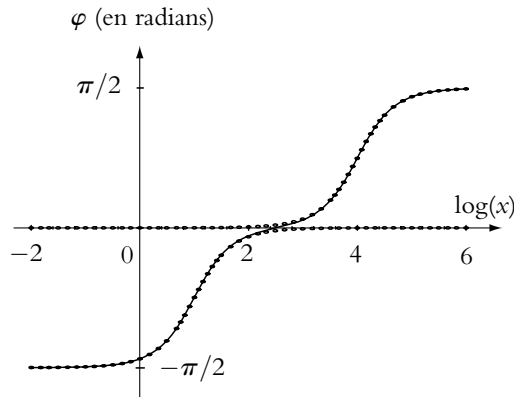


Figure 18.24 Composition de diagrammes de Bode en phase : points : φ_1 et φ_2 , trait plein : φ et $x = \omega/\omega_{\text{ref}}$ avec $\omega_{\text{ref}} = 1 \text{ rad.s}^{-1}$.

6. Filtres du deuxième ordre (PCSI, PTSI)

On s'intéresse aux fonctions de transfert du deuxième ordre, c'est-à-dire celles qui ont comme terme de plus haut degré un terme en ω^2 . Toute l'étude sera menée sur un circuit R, L, C série déjà étudié bien que dans la pratique on utilise très souvent des circuits autres que le circuit R, L, C série mais on pourra toujours se ramener à la même forme canonique.

On utilise les constantes caractéristiques définies dans le chapitre sur l'étude du circuit R, L, C série :

$$\omega_0 = \frac{1}{\sqrt{LC}} \text{ pulsation propre et} \quad (18.17)$$

$$Q = \frac{L\omega_0}{R} = \frac{1}{RC\omega_0} \text{ facteur de qualité} \quad (18.18)$$

Comme pour les filtres d'ordre un, on suppose que la charge a une impédance infinie donc que le courant de sortie du quadripôle est nul.

On utilisera d'autre part les résultats des calculs sur les résonances et sur le circuit R, L, C série effectués dans le chapitre correspondant. Il est donc indispensable que le lecteur maîtrise cette étude.

6.1 Filtre passe-bas du deuxième ordre

Le circuit étudié est celui de la figure 18.25 où la tension de sortie est prise aux bornes du condensateur.

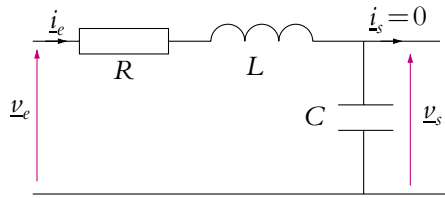


Figure 18.25 Exemple de filtre passe-bas du second ordre.

a) Étude des comportements limites

En T.B.F., le condensateur se comporte comme un circuit ouvert et la bobine comme un court-circuit. On obtient le schéma équivalent de la figure 18.26. L'intensité est nulle dans le circuit donc toute la tension d'entrée se reporte aux bornes de C et $v_s = v_e$ soit $\underline{H} = 1$.

En T.H.F., le condensateur se comporte comme un court-circuit et la bobine comme un circuit ouvert. On obtient le schéma équivalent de la figure 18.27. La tension de sortie est prise aux bornes d'un court-circuit, elle est donc nulle et $\underline{H} = 0$.

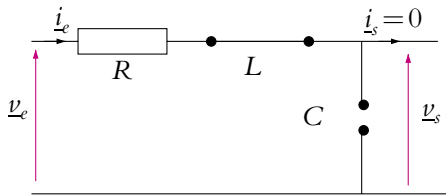


Figure 18.26 Circuit équivalent en basses fréquences.

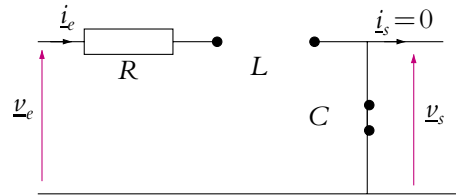


Figure 18.27 Circuit équivalent en hautes fréquences.

Le filtre laisse passer les basses fréquences et coupe les hautes fréquences, il s'agit donc d'un *filtre passe-bas*.

b) Calcul de la fonction de transfert

Comme le courant de sortie est nul, on peut appliquer un diviseur de tension et :

$$\underline{H} = \frac{v_s}{v_e} = \frac{\underline{Z}_c}{R + \underline{Z}_L + \underline{Z}_c} = \frac{1}{\underline{Y}_c R + \underline{Z}_L \underline{Y}_c + 1} = \frac{1}{1 - LC\omega^2 + jRC\omega}$$

On retrouve les comportements limites en faisant tendre ω vers 0 ou vers l'infini. Si l'on utilise les expressions de ω_0 (18.17) et Q (18.18), alors le terme $LC\omega^2$ devient $\left(\frac{\omega}{\omega_0}\right)^2$ et $RC\omega$ devient $\frac{\omega}{\omega_0 Q}$.

On obtient donc la fonction de transfert sous forme canonique :

$$\underline{H} = \frac{1}{1 - \left(\frac{\omega}{\omega_0}\right)^2 + j\frac{\omega}{\omega_0 Q}}$$

En fait, le numérateur n'est pas forcément égal à 1 dans le cas général et il faudra retenir la forme canonique du filtre passe-bas du deuxième ordre sous la forme suivante :

$$\underline{H} = \frac{H_0}{1 - \left(\frac{\omega}{\omega_0}\right)^2 + j\frac{\omega}{\omega_0 Q}} \quad (18.19)$$

On utilise aussi souvent la variable réduite $x = \frac{\omega}{\omega_0}$ et dans ce cas la fonction de transfert s'écrit :

$$\underline{H} = \frac{H_0}{1 - x^2 + j\frac{x}{Q}}$$

► Remarques

- Il faut faire attention à la définition de la variable réduite x . Dans le cas des filtres du premier ordre, c'est la pulsation de coupure qui intervient alors que dans le cas des filtres du second ordre, c'est la pulsation propre.
- Comme pour les filtres du premier ordre, si on réalise un circuit différent du circuit R, L, C série et qu'il se comporte en passe-bas, la forme canonique sera inchangée mais les expressions de H_0 , ω_0 et Q seront différentes de celles du circuit R, L, C série.

c) Étude du gain

Si on utilise la variable réduite x , le gain en décibels est donné par l'expression :

$$G_{dB} = G_0 + 20 \log \frac{1}{\sqrt{(1-x^2)^2 + \left(\frac{x}{Q}\right)^2}}$$

avec $G_0 = 20 \log |H_0|$.

On commence par étudier les comportements asymptotiques. Ici la pulsation de référence est ω_0 .

- Pour $\omega \ll \omega_0$, soit $x \ll 1$, on ne garde que le terme de plus bas degré du dénominateur donc 1. G_{dB} se comporte alors comme son asymptote $G_{dB}^{a_1} = G_0$.
- Pour $\omega \gg \omega_0$, soit $x \gg 1$, on ne garde que le terme de plus haut degré du dénominateur donc $\sqrt{(x^2)^2} = x^2$. G_{dB} se comporte alors comme son asymptote $G_{dB}^{a_2} = G_0 - 20 \log x^2 = G_0 - 40 \log x$. Puisqu'il y a un facteur -40 devant $\log x$, l'asymptote est une droite de pente -40 dB/déc.

Les asymptotes se croisent pour $G_{dB}^{a_1} = G_{dB}^{a_2}$ soit $-40 \log x = 0$ donc pour $x = 1$ ou encore $\omega = \omega_0$.

Il a été vu dans le chapitre sur l'étude du circuit R, L, C série qu'il y a résonance si Q est supérieur à $1/\sqrt{2}$ pour une pulsation ω_r plus petite que ω_0 , $\omega_r = \omega_0 \sqrt{1 - \frac{1}{2Q^2}}$. On obtient alors les courbes de la figure 18.28 tracées pour les trois valeurs de Q : 0,2 ; $1/\sqrt{2}$ et 5 et pour $H_0 = 1$.

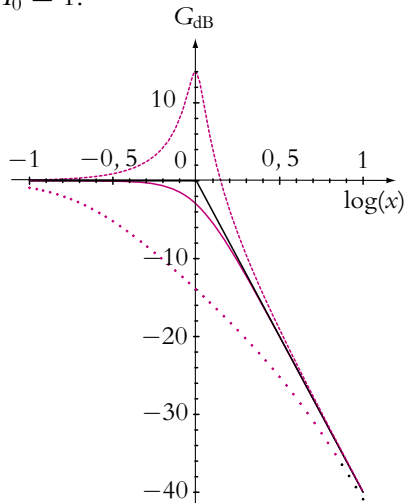


Figure 18.28 Diagramme de Bode en gain d'un filtre passe-bas du second ordre pour différentes valeurs du facteur de qualité : pointillés : $Q = 5$, trait plein : $Q = 1/\sqrt{2}$, croix : $Q = 0,2$.

L'intérêt d'un filtre passe-bas est d'atténuer les hautes fréquences ; or s'il y a résonance, on voit sur la figure que la courbe réelle rejoint moins rapidement l'asymptote. On s'arrange donc en général pour que le facteur de qualité Q soit inférieur à $1/\sqrt{2}$. Dans ce cas, les signaux de hautes fréquences sont mieux atténués qu'avec un filtre du premier ordre puisque la pente de l'asymptote est -40 dB/déc au lieu de -20 dB/déc c'est-à-dire que si la fréquence est multipliée par 10, l'amplitude du signal est divisée par 100 pour un passe-bas du deuxième ordre au lieu de 10 pour un passe-bas du premier ordre.

À partir de cette étude, on peut désigner Q et ω_0 par d'autres termes que « facteur de qualité » et « pulsation propre » : on donne à Q le nom de *facteur de forme*, car c'est de sa valeur que dépend la forme de la courbe et à ω_0 le nom de *facteur d'échelle* car c'est de sa valeur que dépend la position des asymptotes.

La détermination de la bande passante a été faite dans le chapitre sur l'étude du circuit R, L, C série.

d) Étude de la phase

On suppose que $H_0 = 1$. On commence par calculer les équivalents de $\underline{H}$ en T.B.F. et T.H.F..

- En T.B.F., $x \ll 1$ et, en ne gardant que le terme de plus bas degré en x , on trouve $\underline{H} \simeq 1$. Ainsi l'argument de $\underline{H}$ réel positif est $\varphi = 0$.
- En T.H.F., $x \gg 1$ et, en ne gardant que le terme de plus haut degré en x , on trouve $\underline{H} \simeq \frac{-1}{x^2}$. Il y a donc une indétermination puisqu'on ne sait pas si $\varphi = \pi$ ou $\varphi = -\pi$. Il faut donc étudier plus précisément l'intervalle de variation de φ .

Attention au cas où H_0 est différent de 1 : s'il est positif, le résultat n'est pas modifié ; s'il est négatif, il faut ajouter π aux phases déterminées ci-dessus.

En utilisant les mêmes notations que pour les filtres du premier ordre, $\varphi = \varphi_N - \varphi_D$. Or $\varphi_N = 0$ donc :

$$\tan \varphi_D = \frac{x}{Q(1 - x^2)}$$

Ici, il est plus intéressant de regarder le signe de $\sin \varphi_D$ qui est constant alors que le signe du cosinus change. $\sin \varphi_D$ est positif donc $0 < \varphi_D < \pi$ et $-\pi < \varphi < 0$. Cette étude permet de lever l'indétermination de l'étude T.H.F. On retrouve que :

- lorsque $x \rightarrow 0$, $\tan \varphi_D \rightarrow 0$ par valeur positive donc $\varphi_D \rightarrow 0$,
- lorsque $x \rightarrow \infty$, $\tan \varphi_D \rightarrow 0$ par valeur négative donc $\varphi_D \rightarrow \pi$.

On retrouve les valeurs limites de φ déterminées précédemment.

D'autre part, pour $\omega = \omega_0$ ou $x = 1$, $\tan \varphi_D \rightarrow \infty$ donc $\varphi_D = \frac{\pi}{2}$ et $\varphi = -\frac{\pi}{2}$.

On pourrait démontrer par une étude de la dérivée de la fonction $\varphi(x)$ que sa variation en fonction de la pulsation est monotone. On obtient la courbe de la figure 18.29

tracée pour les trois mêmes valeurs de Q que précédemment.

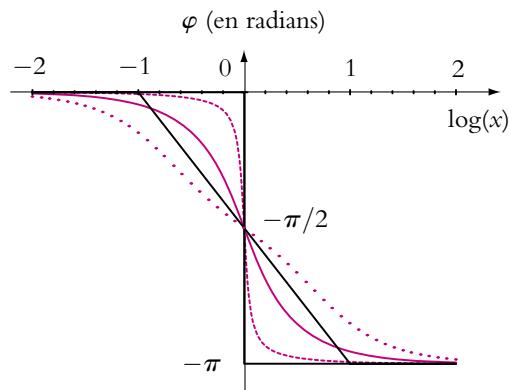


Figure 18.29 Diagramme de Bode en phase d'un filtre passe-bas du second ordre pour différentes valeurs du facteur de qualité : pointillés : $Q = 5$, trait plein : $Q = 1/\sqrt{2}$, croix : $Q = 0, 2$.

6.2 Filtre passe-haut du deuxième ordre

Le circuit étudié est celui de la figure 18.30, où la tension de sortie est prise aux bornes de la bobine. On rappelle qu'expérimentalement on ne peut mesurer la tension aux bornes d'une bobine pure car une bobine réelle comporte toujours une résistance.

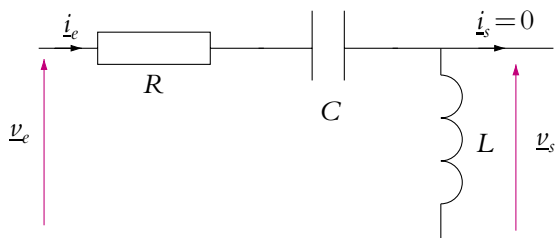


Figure 18.30 Exemple de filtre passe-haut du second ordre.

a) Étude des comportements limites

En T.B.F., la bobine se comporte comme un court-circuit et le condensateur comme un circuit ouvert. On obtient le schéma équivalent de la figure 18.31. La tension de sortie est prise aux bornes d'un court-circuit, elle est donc nulle et $H = 0$.

En T.H.F., la bobine se comporte comme un circuit ouvert et le condensateur comme un court-circuit. On obtient le schéma équivalent de la figure 18.32. L'intensité est nulle dans le circuit donc toute la tension d'entrée se reporte aux bornes de la bobine et $v_s = v_e$ soit $H = 1$.

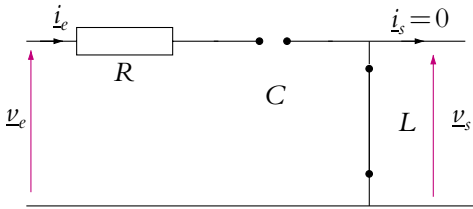


Figure 18.31 Circuit équivalent en basses fréquences.

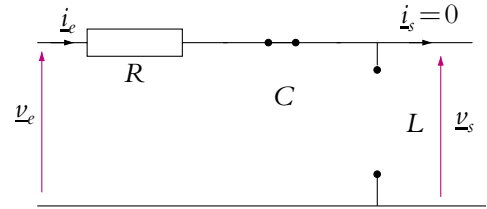


Figure 18.32 Circuit équivalent en hautes fréquences.

Le filtre laisse passer les hautes fréquences et coupe les basses fréquences, il s'agit donc d'un *filtre passe-haut*.

b) Calcul de la fonction de transfert

Comme le courant de sortie est nul, on peut appliquer un diviseur de tension :

$$\underline{H} = \frac{v_s}{v_e} = \frac{\underline{Z}_L}{R + \underline{Z}_L + \underline{Z}_C} = \frac{\underline{Z}_L \underline{Y}_C}{\underline{Y}_C R + \underline{Z}_L \underline{Y}_C + 1} = \frac{-LC\omega^2}{1 - LC\omega^2 + jRC\omega}$$

On retrouve les comportements limites en faisant tendre ω vers 0 ou vers l'infini. Si on utilise les expressions de ω_0 (18.17) et Q (18.18), le terme $LC\omega^2$ devient $\left(\frac{\omega}{\omega_0}\right)^2$ et $RC\omega$ devient $\frac{\omega}{\omega_0 Q}$.

On obtient donc la fonction de transfert sous forme canonique :

$$\underline{H} = \frac{-\left(\frac{\omega}{\omega_0}\right)^2}{1 - \left(\frac{\omega}{\omega_0}\right)^2 + j\frac{\omega}{\omega_0 Q}}$$

En fait, comme pour les autres filtres, pour obtenir la forme canonique générale, il faut multiplier le numérateur par H_0 . On obtient donc l'expression à retenir :

$$\underline{H} = \frac{-\left(\frac{\omega}{\omega_0}\right)^2 H_0}{1 - \left(\frac{\omega}{\omega_0}\right)^2 + j\frac{\omega}{\omega_0 Q}} \quad (18.20)$$

Avec la variable réduite $x = \frac{\omega}{\omega_0}$, la fonction de transfert s'écrit :

$$\underline{H} = \frac{-x^2 H_0}{1 - x^2 + j\frac{x}{Q}}$$

► **Remarque :** Comme pour les autres filtres, si on réalise un circuit différent du R, L, C série et qu'il se comporte en passe-bas, la forme canonique sera inchangée mais les expressions de ω_0 et Q seront différentes de celles du circuit R, L, C série.

c) Étude du gain

Si on utilise la variable réduite x , le gain en décibels est donné par l'expression :

$$G_{dB} = G_0 + 20 \log \frac{x^2}{\sqrt{(1 - x^2)^2 + \left(\frac{x}{Q}\right)^2}}$$

avec $G_0 = 20 \log |H_0|$.

On commence par étudier les comportements asymptotiques. Ici, la pulsation de référence est ω_0 .

- Pour $\omega \ll \omega_0$, soit $x \ll 1$, on ne garde que le terme de plus bas degré du dénominateur donc 1. G_{dB} se comporte alors comme son asymptote :

$$G_{dB}^{a1} = G_0 + 20 \log x^2 = G_0 + 40 \log x$$

Puisqu'il y a un facteur 40 devant $\log x$, l'asymptote est une droite de pente 40 dB/déc.

- Pour $\omega \gg \omega_0$, soit $x \gg 1$, on garde le terme de plus haut degré du dénominateur donc $\sqrt{(x^2)^2} = x^2$. G_{dB} se comporte alors comme son asymptote

$$G_{dB}^{a2} = G_0 - 20 \log \left(\frac{x^2}{x^2} \right) = G_0.$$

Les asymptotes se croisent pour $G_{dB}^{a1} = G_{dB}^{a2}$ soit $40 \log x = 0$ donc $x = 1$ ou encore $\omega = \omega_0$.

Comme pour l'étude de la tension aux bornes du condensateur, on peut établir qu'il y a une résonance si Q est supérieur à $1/\sqrt{2}$ pour une pulsation plus grande que ω_0 , $\omega_r = \omega_0 \frac{1}{\sqrt{1 - \frac{1}{2Q^2}}}$. On obtient alors les

courbes de la figure 18.33 tracées pour les trois valeurs de Q : 0,2 ; $1/\sqrt{2}$ et 5 et pour $H_0 = 1$.

L'intérêt d'un filtre passe-haut est d'atténuer les basses fréquences ; or s'il y a résonance, on voit sur la figure que la courbe réelle rejoint moins rapidement l'asymptote.

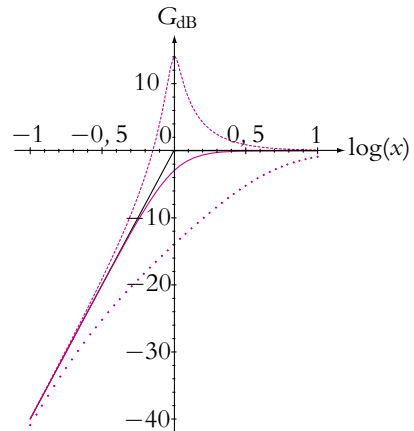


Figure 18.33 Diagramme de Bode en gain d'un filtre passe-haut du second ordre pour différentes valeurs du facteur de qualité : pointillés : $Q = 5$, trait plein : $Q = 1/\sqrt{2}$, croix : $Q = 0,2$.

Comme pour le filtre passe-bas, on s'arrange donc en général pour que le facteur de qualité Q soit inférieur à $1/\sqrt{2}$. Dans ce cas, les signaux de basses fréquences sont mieux atténués qu'avec un filtre du premier ordre puisque la pente de l'asymptote est $+40$ dB/déc au lieu de $+20$ dB/déc ; c'est-à-dire que si la fréquence est divisée par 10, l'amplitude du signal est divisée par 100 pour un passe-haut du deuxième ordre au lieu de 10 pour un passe-haut du premier ordre.

La bande passante se calcule comme celle du filtre passe-bas.

d) Étude de la phase

On suppose dans ce paragraphe que $H_0 = 1$. On commence par calculer les équivalents de $\underline{H}$ en T.B.F. et T.H.F.

- En T.B.F., $x \ll 1$ et, en ne gardant que le terme de plus bas degré en x au dénominateur, on trouve $\underline{H} \simeq -x^2$. Ainsi l'argument de $\underline{H}$, réel négatif, est $\varphi = -\pi$ ou $\varphi = \pi$.
- En T.H.F., $x \gg 1$ et, en ne gardant que le terme de plus haut degré en x au dénominateur, on trouve $\underline{H} \simeq \frac{-x^2}{-x^2}$ donc $\varphi = 0$.

Attention au cas où H_0 est différent de 1 : s'il est positif, le résultat n'est pas modifié, s'il est négatif, il faut ajouter π aux phases déterminées ci-dessus.

En utilisant les mêmes notation que précédemment, $\varphi = \varphi_N - \varphi_D$ avec $\varphi_N = \pi$ et :

$$\tan \varphi_D = \frac{x}{Q(1 - x^2)}$$

L'étude de φ_D est la même que pour le filtre passe-bas du deuxième ordre. $\sin \varphi_D$ est positif donc $0 < \varphi_D < \pi$ et $0 < \varphi < \pi$.

- Lorsque $x \rightarrow 0$, $\tan \varphi_D \rightarrow 0$ par valeur positive donc $\varphi_D \rightarrow 0$ et $\varphi = \pi - \varphi_D \rightarrow \pi$.
- Lorsque $x \rightarrow \infty$, $\tan \varphi_D \rightarrow 0$ par valeur négative donc $\varphi_D \rightarrow \pi$ et $\varphi \rightarrow 0$.

On retrouve les valeurs limites de φ déterminées précédemment.

D'autre part, pour $\omega = \omega_0$ ou $x = 1$, $\tan \varphi_D \rightarrow \infty$ donc $\varphi_D = \frac{\pi}{2}$ et $\varphi = \frac{\pi}{2}$.

Comme pour le filtre passe-bas, on pourrait démontrer par une étude de la dérivée de la fonction φ que sa variation en fonction de la pulsation est monotone. On obtient la courbe de la figure 18.34 tracée pour trois valeurs de Q : 0,2 ; $1/\sqrt{2}$ et 5 et pour $H_0 = 1$.

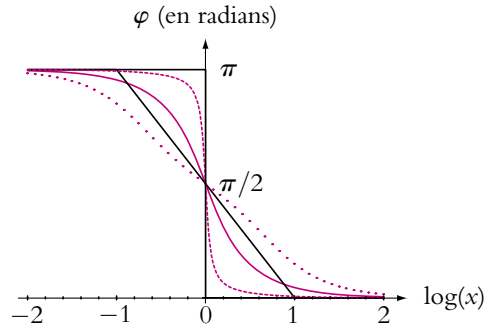


Figure 18.34 Diagramme de Bode en phase d'un filtre passe-haut du second ordre pour différentes valeurs du facteur de qualité : pointillés : $Q = 5$, trait plein : $Q = 1/\sqrt{2}$, croix : $Q = 0,2$.

6.3 Filtre passe-bande du deuxième ordre

Le circuit étudié est celui de la figure 18.35 où la tension de sortie est prise aux bornes du résistor.

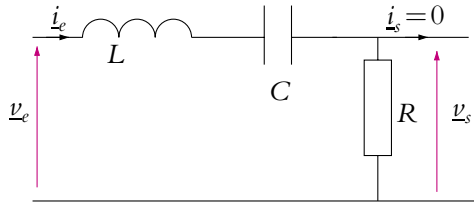


Figure 18.35 Exemple de filtre passe-bande du second ordre.

a) Étude des comportements limites

En T.B.F., la bobine se comporte comme un court-circuit et le condensateur comme un circuit ouvert. On obtient le schéma équivalent de la figure 18.36. Le courant dans R est nul donc $v_s = 0$ V et $H = 0$.

En T.H.F., la bobine se comporte comme un circuit ouvert et le condensateur comme un court-circuit. On obtient le schéma équivalent de la figure 18.37. Le courant dans R est nul donc $v_s = 0$ V et $H = 0$.

D'après l'étude faite dans le chapitre sur le circuit R, L, C série, on a vu que l'impédance du circuit R, L, C est minimale pour $\omega = \omega_0$. Pour une tension d'entrée donnée, l'amplitude de l'intensité sera donc maximale pour $\omega = \omega_0$ et donc l'amplitude de v_s aussi.

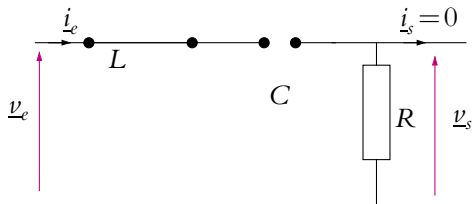


Figure 18.36 Circuit équivalent en basses fréquences.

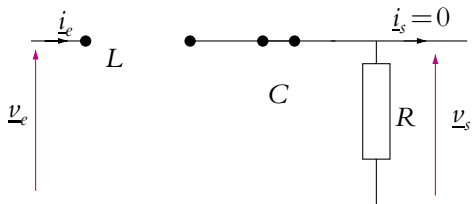


Figure 18.37 Circuit équivalent en hautes fréquences.

Le filtre laisse passer les fréquences intermédiaires mais coupe les hautes fréquences et les basses fréquences, il s'agit donc d'un *filtre passe-bande*.

b) Calcul de la fonction de transfert

Comme le courant de sortie est nul, on peut appliquer un diviseur de tension :

$$\underline{H} = \frac{v_s}{v_e} = \frac{R}{R + \underline{Z}_L + \underline{Z}_C}$$

Dans le cas de ce filtre, deux formes sont utiles. Il suffit d'en retenir une mais il faut savoir que les deux existent. En mettant R en facteur au numérateur et au dénominateur, on obtient :

$$\underline{H} = \frac{1}{1 + j \left(\frac{L\omega}{R} - \frac{1}{RC\omega} \right)}$$

Cette forme est utile pour trouver sans calcul la pulsation de résonance du filtre comme on le verra dans la suite. En multipliant numérateur et dénominateur par $\underline{Y}_c$:

$$\underline{H} = \frac{\underline{Y}_c R}{1 + \underline{Y}_c R + \underline{Z}_L \underline{Y}_c} = \frac{jRC\omega}{1 - LC\omega^2 + jRC\omega}$$

L'intérêt de cette forme est que le dénominateur est le même que pour les deux précédents filtres étudiés, elle est donc plus facile à retenir.

On retrouve les comportements limites en faisant tendre ω vers 0 ou vers l'infini. Si on utilise les expressions de ω_0 (18.17) et Q (18.18), le terme $LC\omega^2$ devient $\left(\frac{\omega}{\omega_0}\right)^2$ et $RC\omega$ devient $\frac{\omega}{\omega_0 Q}$.

On obtient donc la fonction de transfert sous deux formes canoniques :

$$\underline{H} = \frac{1}{1 + jQ \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)}$$

ou

$$\underline{H} = \frac{\frac{j\omega}{\omega_0 Q}}{1 - \left(\frac{\omega}{\omega_0}\right)^2 + j\frac{\omega}{\omega_0 Q}}$$

En fait, comme pour les autres filtres, pour obtenir les formes canoniques générales, il faut multiplier le numérateur par H_0 . On obtient donc les expressions :

$$\underline{H} = \frac{H_0}{1 + jQ \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)} \quad (18.21)$$

ou

$$\underline{H} = \frac{\frac{j\omega}{\omega_0 Q} H_0}{1 - \left(\frac{\omega}{\omega_0}\right)^2 + j\frac{\omega}{\omega_0 Q}} \quad (18.22)$$

Avec la variable réduite $x = \frac{\omega}{\omega_0}$, la fonction de transfert s'écrit :

$$\underline{H} = \frac{H_0}{1 + jQ \left(x - \frac{1}{x} \right)} = \frac{j \frac{x}{Q} H_0}{1 - x^2 + j \frac{x}{Q}}$$

On passe de la deuxième forme à la première en divisant par $\frac{jx}{Q}$.

► **Remarque :** Comme pour les autres filtres, si on réalise un circuit différent du circuit R, L, C série et s'il se comporte en passe-bande, la forme canonique sera inchangée mais les expressions de H_0 , ω_0 et Q seront différentes de celles du circuit R, L, C série.

Pour trouver la pulsation de résonance, on utilise la forme (18.21) et il n'y a pas besoin de calculer de dérivée. Le module de $\underline{H}$ est :

$$H = \frac{|H_0|}{\sqrt{1 + Q^2 \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)^2}}$$

H est maximal quand le dénominateur est minimal. Or le dénominateur est la somme de deux termes positifs dont l'un est constant et l'autre positif ou nul, il est donc minimal lorsque le deuxième terme est lui aussi minimal donc nul soit pour $\omega = \omega_0$. On retrouve sans calcul que ω_0 est la pulsation de résonance.

c) Étude du gain

Si on utilise la variable réduite x , le gain en décibels est donné par l'expression :

$$G_{dB} = G_0 + 20 \log \frac{\frac{x}{Q}}{\sqrt{(1 - x^2)^2 + \left(\frac{x}{Q} \right)^2}}$$

avec $G_0 = 20 \log |H_0|$, ou

$$G_{dB} = G_0 - 20 \log \sqrt{1 + Q^2 \left(x - \frac{1}{x} \right)^2}$$

On commence par étudier les comportements asymptotiques sur la première forme avec comme pulsation de référence ω_0 . Le lecteur peut s'entraîner sur la deuxième forme et vérifier qu'il trouve les mêmes résultats.

- Pour $\omega \ll \omega_0$, donc pour $x \ll 1$, on ne garde que le terme de plus bas degré du dénominateur donc 1. G_{dB} se comporte alors comme son asymptote :

$$G_{dB}^{a1} = G_0 + 20 \log \frac{x}{Q} = G_0 + 20 \log x - 20 \log Q$$

Puisque il y a un facteur 20 devant $\log x$, l'asymptote est une droite de pente 20 dB/déc.

- Pour $\omega \gg \omega_0$, donc pour $x \gg 1$, on ne garde que le terme de plus haut degré du dénominateur donc $\sqrt{(x^2)^2} = x^2$. G_{dB} se comporte alors comme son asymptote :

$$G_{dB}^{a_2} = G_0 + 20 \log \frac{x}{Qx^2} = G_0 - 20 \log x - 20 \log Q$$

Puisque il y a un facteur -20 devant $\log x$ l'asymptote est une droite de pente -20 dB/déc.

Les asymptotes se croisent pour $G_{dB}^{a_1} = G_{dB}^{a_2}$, soit , $20 \log \frac{x}{Q} = 20 \log \frac{1}{xQ}$ donc $x = 1$ ou encore $\omega = \omega_0$. L'ordonnée des asymptotes à leur croisement est $G_{dB}^{a_1}(x = 1) = G_0 - 20 \log Q$ et dépend du facteur de qualité. S'il est supérieur à 1, le croisement s'effectue pour une ordonnée inférieure à G_0 , et si Q est inférieur à 1, le croisement s'effectue pour une ordonnée supérieure à G_0 .

On a vu au paragraphe précédent que la résonance a lieu pour $x = 1$ et qu'alors $\underline{H} = H_0$ donc $G_{dB}(x = 1) = G_0$. Par conséquent, quel que soit le facteur de qualité, le maximum est le même.

On obtient alors les courbes de la figure 18.38 tracées pour les trois valeurs de Q : 0,2 ; $1/\sqrt{2}$ et 5 et $H_0 = 1$.

On s'aperçoit que pour un facteur de qualité élevé, la courbe est beaucoup plus resserrée autour de ω_0 ($x = 1$). Le filtre est dans ce cas dit *sélectif*, c'est-à-dire qu'il laisse passer les signaux de pulsation proche de ω_0 mais atténue les autres. La largeur de la bande passante a été déterminée dans le chapitre sur l'étude des circuits R, L, C série, on rappelle son expression :

$$\Delta\omega = \frac{\omega_0}{Q}$$

Cette expression confirme la conclusion précédente tirée de l'allure de la courbe : plus Q est grand, plus la bande passante est étroite.

► **Remarque** : On voit sur la figure 18.38 et on peut établir par le calcul que la courbe réelle est au-dessus de ses asymptotes pour $Q > 1$ et en-dessous pour $Q \leq 1$.

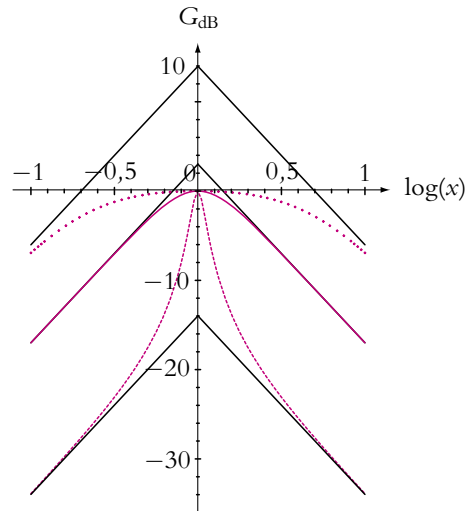


Figure 18.38 Diagramme de Bode en gain d'un filtre passe-bande du second ordre pour différentes valeurs du facteur de qualité : pointillés : $Q = 5$, trait plein : $Q = 1/\sqrt{2}$, croix : $Q = 0,2$.

d) Étude de la phase

On fait l'étude dans le cas $H_0 = 1$. On commence par calculer les équivalents de $\underline{H}$ en T.B.F. et T.H.F.

- En T.B.F., $x \ll 1$ et, en ne gardant que le terme de plus bas degré en x au dénominateur, on trouve $\underline{H} \simeq \frac{jx}{Q}$. Ainsi l'argument de $\underline{H}$ est $\varphi = \frac{\pi}{2}$.
- En T.H.F., $x \gg 1$ et, en ne gardant que le terme de plus haut degré en x au dénominateur, on trouve $\underline{H} \simeq \frac{jx}{-x^2 Q}$ donc $\varphi = -\frac{\pi}{2}$.

Attention comme pour les filtres précédents au cas où H_0 est différent de 1 : s'il est positif, le résultat n'est pas modifié ; s'il est négatif, il faut ajouter π aux phases déterminées ci-dessus.

Il semble donc que la phase soit comprise entre $\pi/2$ et $-\pi/2$, ce que va confirmer l'étude complète. En utilisant les mêmes notations que précédemment, $\varphi = \varphi_N - \varphi_D$ avec $\varphi_N = \frac{\pi}{2}$ et :

$$\tan \varphi_D = \frac{x}{Q(1 - x^2)}$$

L'étude de φ_D est encore la même que pour le filtre passe-bas du deuxième ordre. $\sin \varphi_D$ est positif donc $0 < \varphi_D < \pi$ et $-\frac{\pi}{2} < \varphi < \frac{\pi}{2}$.

- Lorsque $x \rightarrow 0$, $\tan \varphi_D \rightarrow 0$ par valeur positive donc $\varphi_D \rightarrow 0$ et $\varphi = \frac{\pi}{2} - \varphi_D \rightarrow \frac{\pi}{2}$.
- Lorsque $x \rightarrow \infty$, $\tan \varphi_D \rightarrow 0$ par valeur négative donc $\varphi_D \rightarrow \pi$ et $\varphi \rightarrow -\frac{\pi}{2}$.

On retrouve les valeurs limites de φ déterminées précédemment.

D'autre part, pour $\omega = \omega_0$ ou $x = 1$, $\tan \varphi_D \rightarrow \infty$ donc $\varphi_D = \frac{\pi}{2}$ et $\varphi = 0$.

Le déphasage est nul à la résonance du filtre passe-bande.

Dans le cas général, si H_0 est négatif, le déphasage est égal à π à la résonance.

Comme pour le filtre passe-bas, on pourrait démontrer par une étude de la dérivée de la fonction φ que sa variation en fonction de la pulsation est monotone. On obtient la courbe de la figure 18.39 tracée pour les trois mêmes valeurs de Q que précédemment.

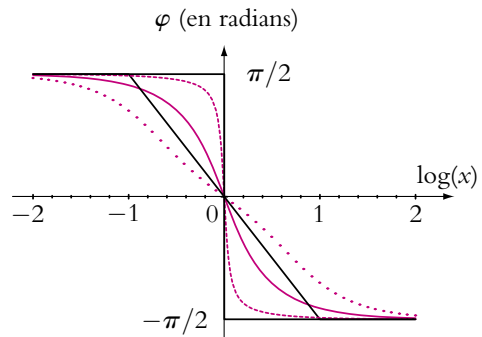


Figure 18.39 Diagramme de Bode en phase d'un filtre passe-bande du second ordre pour différentes valeurs du facteur de qualité : $Q = 5$, trait pointillés ; $Q = 1/\sqrt{2}$, croix ; $Q = 0,2$.

Aux limites de la bande passante $\varphi = \pm \frac{\pi}{4}$.

e) Comportements intégrateur et dérivateur du filtre passe-bande

En basses fréquences, la courbe de gain est confondue avec une asymptote de pente +20 dB/déc donc le filtre a un comportement dérivateur comme le filtre passe-haut du premier ordre.

En hautes fréquences, la courbe de gain est confondue avec une asymptote de pente -20 dB/déc donc le filtre a un comportement intégrateur comme le filtre passe-bas du premier ordre.

6.4 Filtre coupe-bande ou réjecteur de bande du deuxième ordre

Le circuit étudié est celui de la figure 18.40, où la tension de sortie est prise aux bornes de l'ensemble condensateur et bobine.

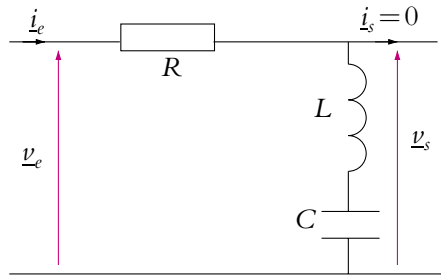


Figure 18.40 Exemple de filtre coupe-bande du second ordre.

a) Étude des comportements limites

En T.B.F., la bobine se comporte comme un court-circuit et le condensateur comme un circuit ouvert. On obtient le schéma équivalent de la figure 18.41. Le courant dans R est nul donc $v_s = v_e$ et $H = 1$.

En T.H.F., la bobine se comporte comme un circuit ouvert et le condensateur comme un court-circuit. On obtient le schéma équivalent de la figure 18.42. Le courant dans R est nul donc $v_s = v_e$ et $H = 1$.

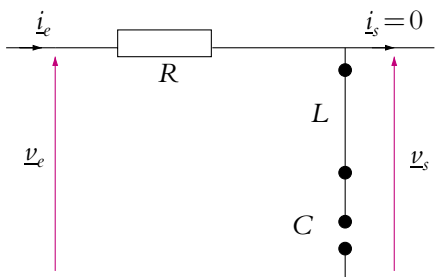


Figure 18.41 Circuit équivalent en basses fréquences.

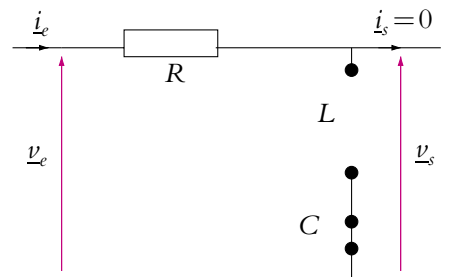


Figure 18.42 Circuit équivalent en hautes fréquences.

D'autre part, l'impédance du dipôle formé de L et C est $\underline{Z} = j \left(L\omega - \frac{1}{C\omega} \right)$, elle est nulle pour $\omega = \omega_0$ et donc $v_s(\omega_0) = 0$.

Le filtre coupe les fréquences intermédiaires mais laisse passer les hautes fréquences et les basses fréquences, il s'agit d'un *filtre coupe-bande ou réjecteur de bande*.

b) Calcul de la fonction de transfert

Comme le courant de sortie est nul, on peut appliquer un diviseur de tension et :

$$\underline{H} = \frac{v_s}{v_e} = \frac{\underline{Z}_L + \underline{Z}_C}{R + \underline{Z}_L + \underline{Z}_C}$$

En multipliant numérateur et dénominateur par $\underline{Y}_C$:

$$\underline{H} = \frac{1 + \underline{Z}_L \underline{Y}_C}{1 + \underline{Y}_C R + \underline{Z}_L \underline{Y}_C} = \frac{1 - LC\omega^2}{1 - LC\omega^2 + jRC\omega}$$

On retrouve encore une fois le même dénominateur.

On retrouve les comportements limites en faisant tendre ω vers 0 ou vers l'infini. Si on utilise les expressions de ω_0 (18.17) et Q (18.18), le terme $LC\omega^2$ devient $\left(\frac{\omega}{\omega_0} \right)^2$ et $RC\omega$ devient $\frac{\omega}{\omega_0 Q}$.

On obtient donc la fonction de transfert sous forme canonique :

$$\underline{H} = \frac{1 - \left(\frac{\omega}{\omega_0} \right)^2}{1 - \left(\frac{\omega}{\omega_0} \right)^2 + j \frac{\omega}{\omega_0 Q}}$$

En fait, comme pour les autres filtres, pour obtenir les formes canoniques générales, il faut multiplier le numérateur par H_0 . On obtient donc l'expression :

$$\underline{H} = \frac{H_0 \left(1 - \left(\frac{\omega}{\omega_0} \right)^2 \right)}{1 - \left(\frac{\omega}{\omega_0} \right)^2 + j \frac{\omega}{\omega_0 Q}} \quad (18.23)$$

Avec la variable réduite $x = \frac{\omega}{\omega_0}$, la fonction de transfert s'écrit :

$$\underline{H} = \frac{H_0(1 - x^2)}{1 - x^2 + j \frac{x}{Q}}$$

► **Remarques** Comme pour les autres filtres, si on réalise un circuit différent du circuit R, L, C série et s'il se comporte en coupe-bande, la forme canonique sera inchangée mais les expressions de H_0 , ω_0 et Q seront différentes de celles du circuit R, L, C série.

Pour trouver la pulsation rejetée, il suffit d'annuler le numérateur ($H = 0$), en vérifiant que le dénominateur ne s'annule pas pour la même pulsation. On retrouve que ω_0 est la pulsation rejetée.

c) Étude du gain

Si on utilise la variable réduite x , le gain en décibels est donné par l'expression :

$$G_{dB} = 20 \log \frac{|H_0| |1 - x^2|}{\sqrt{(1 - x^2)^2 + \left(\frac{x}{Q}\right)^2}}$$



Attention à la valeur absolue au numérateur.

On commence par étudier les comportements asymptotiques avec comme pulsation de référence ω_0 .

On pose dans la suite $G_0 = 20 \log |H_0|$.

- Pour $\omega \ll \omega_0$, donc pour $x \ll 1$, on ne garde que les termes de plus bas degré du dénominateur et du numérateur donc 1. G_{dB} se comporte alors comme son asymptote $G_{dB}^{a1} = G_0$.
- Pour $\omega \gg \omega_0$ donc pour $x \gg 1$, on ne garde que les termes de plus haut degré du dénominateur et du numérateur donc x^2 et $\sqrt{(x^2)^2} = x^2$. G_{dB} se comporte alors comme son asymptote $G_{dB}^{a2} = G_0 + 20 \log(x^2/x^2) = G_0$.

On obtient donc deux asymptotes horizontales en T.B.F. et T.H.F. D'autre part, pour $\omega = \omega_0$, $H = 0$ donc $G_{dB}(\omega_0) \rightarrow -\infty$. On a donc une troisième asymptote verticale pour $\omega = \omega_0$.

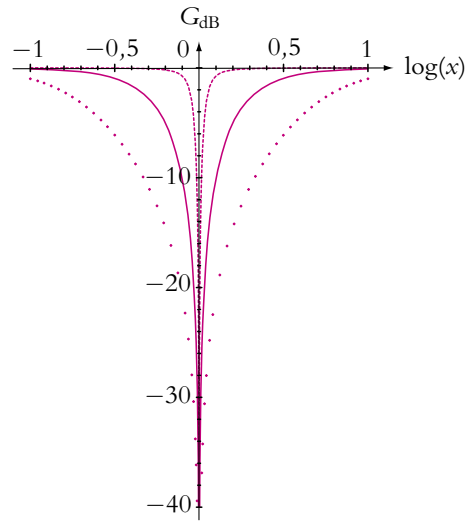


Figure 18.43 Diagramme de Bode en gain d'un filtre coupe-bande du second ordre pour différentes valeurs du facteur de qualité : pointillés : $Q = 5$, trait plein : $Q = 1/\sqrt{2}$, croix : $Q = 0,2$.

On obtient alors les courbes de la figure 18.43 tracées pour les trois valeurs de Q : $0,2$; $1/\sqrt{2}$ et 5 et $H_0 = 1$.

On s'aperçoit que, pour un facteur de qualité élevé, la courbe est beaucoup plus resserrée autour de ω_0 ($x = 1$). Le filtre rejette les signaux de pulsation proche de ω_0 mais laisse passer les autres.

d) Bande rejetée

À l'image de la bande passante déterminée pour un filtre passe-bande, on définit la bande rejetée à -3 dB comme l'intervalle de pulsation dans lequel :

$$H(\omega) \leq \frac{H_{\max}}{\sqrt{2}}$$

ou

$$G_{\text{dB}}(\omega) \leq G_{\text{dB,max}} - 3 \text{ dB}$$

Ici $H_{\max} = |H_0|$ donc on cherche les deux pulsations (appelées aussi pulsations de coupure) telles que $H(\omega) = 1/\sqrt{2}$. En utilisant la variable réduite x :

$$\frac{|H_0||1 - x^2|}{\sqrt{(1 - x^2)^2 + \left(\frac{x}{Q}\right)^2}} = \frac{|H_0|}{\sqrt{2}}$$

soit en élevant au carré et en effectuant les produits en croix :

$$2(1 - x^2)^2 = (1 - x^2)^2 + \frac{x^2}{Q^2} \Leftrightarrow (1 - x^2)^2 = \frac{x^2}{Q^2}$$

À cette étape du calcul, il ne faut absolument pas développer les expressions précédentes mais prendre la racine carrée ou factoriser. On obtient alors deux équations du second degré à résoudre :

$$(a) \quad 1 - x^2 = \frac{x}{Q},$$

$$(b) \quad 1 - x^2 = -\frac{x}{Q}.$$

La première admet comme solutions :

$$x = \frac{\omega}{\omega_0} = -\frac{1}{2Q} \pm \frac{1}{2}\sqrt{\frac{1}{Q^2} + 4}$$

et la deuxième :

$$x = \frac{\omega}{\omega_0} = +\frac{1}{2Q} \pm \frac{1}{2}\sqrt{\frac{1}{Q^2} + 4}$$

En ne gardant que les deux solutions positives, on trouve les deux pulsations de coupure (identiques à celles du filtre passe-bande) :

$$\omega_b = \omega_0 \left(-\frac{1}{2Q} + \frac{1}{2}\sqrt{\frac{1}{Q^2} + 4} \right) \quad \text{et} \quad \omega_h = \omega_0 \left(\frac{1}{2Q} + \frac{1}{2}\sqrt{\frac{1}{Q^2} + 4} \right) \quad (18.24)$$

La largeur de la bande rejetée est obtenue par différence des deux pulsations de coupure :

$$\Delta\omega = \omega_h - \omega_b = \frac{\omega_0}{Q} \quad (18.25)$$

C'est la même expression que pour la bande passante et on retrouve que la bande rejetée est d'autant plus étroite que le facteur de qualité est grand.

e) Étude de la phase

Comme pour les autres filtres, l'étude sera faite pour $H_0 = 1$. On commence par calculer les équivalents de $\underline{H}$ en T.B.F. et T.H.F.

- En T.B.F., $x \ll 1$ et, en ne gardant que les termes de plus bas degré en x au dénominateur et au numérateur, on trouve $\underline{H} \simeq 1$. Ainsi l'argument de $\underline{H}$, réel positif, est $\varphi = 0$.
- En T.H.F., $x \gg 1$ et, en ne gardant que les termes de plus haut degré en x au dénominateur et au numérateur, on trouve $\underline{H} \simeq \frac{x^2}{x^2}$ donc $\varphi = 0$.

Attention comme pour les filtres précédents au cas où H_0 est différent de 1 : s'il est positif, le résultat n'est pas modifié ; s'il est négatif, il faut ajouter π aux phases déterminées ci-dessus.

En utilisant toujours les mêmes notations, $\varphi = \varphi_N - \varphi_D$. On s'aperçoit alors qu'il n'est pas facile de faire l'étude car φ_N et φ_D dépendent de la pulsation. On a intérêt à diviser numérateur et dénominateur par $1 - x^2$, ce qui donne :

$$\underline{H} = \frac{1}{1 + \frac{jx}{Q(1-x^2)}}$$

alors $\varphi_N = 0$ et on peut faire l'étude de φ_D plus facilement.

$$\tan \varphi_D = \frac{x}{Q(1-x^2)}$$

Ici, il est plus intéressant de regarder le signe de $\cos \varphi_D$ qui est positif quel que soit x (alors que $\sin \varphi_D$ change de signe) donc $-\frac{\pi}{2} < \varphi_D < \frac{\pi}{2}$ et $-\frac{\pi}{2} < \varphi < \frac{\pi}{2}$.

- Lorsque $x \rightarrow 0$, $\tan \varphi_D \rightarrow 0$ donc $\varphi_D \rightarrow 0$ et $\varphi = -\varphi_D \rightarrow 0$.
- Lorsque $x \rightarrow \infty$, $\tan \varphi_D \rightarrow 0$ donc $\varphi_D \rightarrow 0$ et $\varphi \rightarrow 0$.

On retrouve les valeurs limites de φ déterminées précédemment. Il est intéressant de s'occuper de ce qui se passe en $x = 1$. Pour $x \rightarrow 1^-$, $\tan \varphi_D \rightarrow +\infty$ donc $\varphi_D \rightarrow \frac{\pi}{2}$ et $\varphi \rightarrow -\frac{\pi}{2}$. Pour $x \rightarrow 1^+$, $\tan \varphi_D \rightarrow -\infty$ donc $\varphi_D \rightarrow -\frac{\pi}{2}$ et $\varphi \rightarrow \frac{\pi}{2}$.

On obtient la courbe de la figure 18.44 tracée pour les trois mêmes valeurs de Q que précédemment.

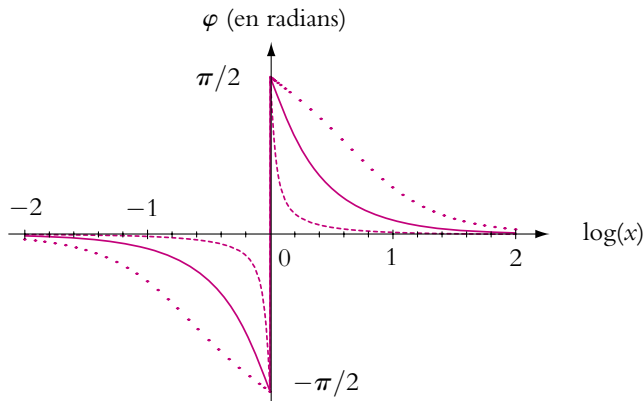


Figure 18.44 Diagramme de Bode en phase d'un filtre coupe-bande du second ordre pour différentes valeurs du facteur de qualité : pointillés : $Q = 5$, trait plein : $Q = 1/\sqrt{2}$, croix : $Q = 0,2$.

► **Remarque** : La discontinuité de phase en $\omega = \omega_0$ n'a pas de réalité physique. Cela signifierait sinon qu'en tournant régulièrement le potentiomètre de réglage de la fréquence, on verrait la phase « faire un bond » de π au passage par ω_0 . Cette discontinuité non physique est due au fait que la bobine a été considérée parfaite et que sa résistance a été négligée. En fait, dans la réalité, la phase varie très rapidement au voisinage de ω_0 mais ne subit pas de discontinuité.

6.5 Récapitulatif sur les filtres du deuxième ordre

Les pentes des asymptotes sont au maximum de ± 40 dB/déc et la différence des pentes est aussi au maximum de $40 \text{ dB} = 2 \times 20 \text{ dB}$.

La variation maximale de la phase est $\pi = 2 \times \frac{\pi}{2}$.

Que ce soit pour la phase ou pour les pentes des asymptotes, les variations trouvées sont deux fois plus grandes pour les filtres d'ordre 2 que pour les filtres d'ordre 1.

Tout polynôme à valeurs complexes de degré supérieur à 2 est factorisable en polynômes de degré 1 ou 2. Ainsi toute fonction de transfert est décomposable en fonction du premier et du second ordre. C'est pour cette raison que l'on se limite à l'étude des filtres du premier et du deuxième ordre.

7. Lien entre fonction de transfert et équation différentielle. Transformations réciproques

On retrouve l'ensemble des termes de la fonction de transfert dans l'équation différentielle lorsque le second membre est une excitation sinusoïdale. Il s'agit donc de deux objets pour décrire une même réalité : le comportement d'un système soumis à un signal sinusoïdal.

La fonction de transfert peut être obtenue soit à partir de l'équation différentielle en utilisant l'analogie entre dérivation dans l'espace des réels et multiplication par $j\omega$ dans l'espace des complexes, soit directement en écrivant des lois des nœuds et des lois des mailles en notation complexe.

Le passage de la fonction de transfert à l'équation différentielle est également possible. Cette approche est particulièrement importante car elle fournit une nouvelle façon d'obtenir l'équation différentielle relative à un système :

- établir la fonction de transfert du système,
- remplacer formellement $j\omega$ par une dérivation par rapport au temps.

Cette deuxième solution est souvent plus rapide, notamment lorsqu'on a un condensateur ou une bobine pour lesquels on ne sait pas forcément *a priori* quelle variable choisir entre tension et intensité. Par exemple, il est difficile d'appliquer la loi des mailles avec un condensateur puisque les tensions s'ajoutent et que la relation fait intervenir l'intensité et la dérivée de la tension par rapport au temps. Pour les mêmes raisons, en inversant tension et intensité, il n'est pas simple d'appliquer la loi des nœuds pour une bobine pour laquelle les intensités s'ajoutent et pour laquelle on a une relation entre la tension et la dérivée de l'intensité. Dans de telles situations, on préférera appliquer la notion de fonction de transfert puis la traduire en équation différentielle.

On pourra le faire que l'entrée soit ou non sinusoïdale. Plusieurs arguments peuvent le justifier. Le premier consiste à considérer cette approche comme une méthode purement formelle basée sur la correspondance parfaite entre dérivation réelle et multiplication par $j\omega$ dans l'espace des complexes. Par ailleurs, dans le chapitre suivant, on verra qu'un signal périodique peut être décomposé sous la forme d'une somme de sinusoïdes. Pour chaque terme de la décomposition, le passage par la fonction de transfert est parfaitement justifié. Les équations étant linéaires, on peut alors sommer les contributions des différentes sinusoïdes pour obtenir l'équation finale. On notera que cette justification pourra se généraliser au cas de signaux non périodiques suffisamment réguliers en utilisant la transformation de Fourier. Il sera également vu dans le cadre du cours de Sciences Industrielles et dans celui du cours de seconde année option PSI une autre généralisation avec la transformée de Laplace.



Il ne faut cependant pas oublier que la notion de fonction de transfert a été introduite dans le cadre du régime sinusoïdal forcé, aspect fortement rappelé par la présence des termes en $j\omega$. Il faudra donc être très prudent dans l'utilisation des fonctions de transfert pour obtenir les équations différentielles afin de ne pas raconter de grosses bêtises.

On retiendra donc que les opérations de multiplication et de division par $j\omega$ au niveau de la fonction de transfert se traduisent respectivement par des dérivations et des intégrations pour les équations différentielles.

Exemple : équation différentielle du filtre de Wien

On considère le montage suivant de la figure 18.45.

Il s'agit du filtre de Wien dont on cherche à établir l'équation différentielle en v , u étant une tension imposée.

La recherche de l'équation différentielle à partir des lois des nœuds et des mailles est délicate à obtenir. On va donc passer par la fonction de transfert.

Celle-ci s'obtient par un pont diviseur de tension :

$$\frac{v}{u} = \frac{Z_{\parallel}}{Z_{\parallel} + Z_s}$$

en notant $Z_{\parallel}$ l'impédance de l'association en parallèle de R et de C et Z_s l'impédance de l'association en série de R et de C .

Ces impédances vérifient :

$$\frac{1}{Z_{\parallel}} = \frac{1}{R} + jC\omega$$

et

$$Z_s = R + \frac{1}{jC\omega}$$

On en déduit :

$$\begin{aligned} \frac{v}{u} &= \frac{1}{1 + \frac{Z_s}{Z_{\parallel}}} \\ &= \frac{1}{1 + \left(\frac{1}{R} + jC\omega\right) \left(R + \frac{1}{jC\omega}\right)} \\ &= \frac{1}{3 + jRC\omega + \frac{1}{jRC\omega}} \\ &= \frac{jRC\omega}{1 + 3RC\omega + (RC)^2(j\omega)^2} \end{aligned}$$

soit l'équation différentielle :

$$(RC)^2 \frac{d^2v}{dt^2} + 3RC \frac{dv}{dt} + v = RC \frac{du}{dt} \quad (8.1)$$

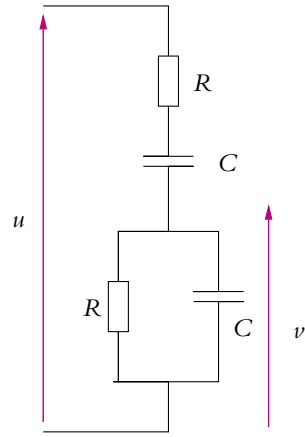


Figure 18.45 Filtre de Wien.

Plusieurs exercices étudient des circuits comportant un amplificateur opérationnel. Le lecteur est donc invité à se reporter à ce chapitre avant de les résoudre.

A. Applications directes du cours

1. Fonction de transfert d'un pont, d'après ENAC 2005

On considère le circuit représenté sur le schéma de la figure (18.46). Un pont dont les quatre branches sont constituées par trois résistors et un condensateur est alimenté par une source de tension sinusoïdale $v_e(t) = v_C - v_E = V_{E0} \cos(\omega t)$, de pulsation ω , connectée aux bornes de la diagonale ED . On désigne par $v_s(t) = v_A - v_B = V_{s0} \cos(\omega t + \varphi_1)$ la tension de sortie recueillie aux bornes de la diagonale AB .

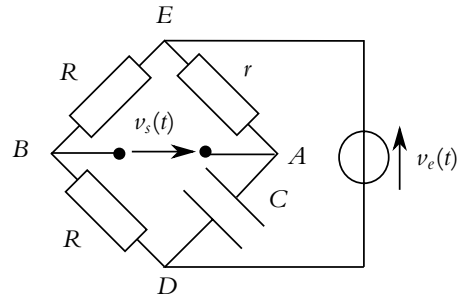


Figure 18.46

On définit la fonction de transfert $\underline{T}(j\omega)$ du circuit par le rapport de l'amplitude complexe $\underline{V_s}$ associée à la tension de sortie sur l'amplitude complexe $\underline{V_e}$, associée à la tension d'entrée.

1. Exprimer $\underline{T}(j\omega) = \underline{V_s}/\underline{V_e}$.
2. Calculer le module de $\underline{T}$. Exprimer le déphasage φ_1 de la tension de sortie $v_s(t)$ par rapport à la tension d'entrée $v_e(t)$. Tracer $\varphi_1(\omega)$ en fonction de ω .
3. On donne $\omega = 1000 \text{ rad.s}^{-1}$, $C = 1 \text{ } \mu\text{F}$. Quelle valeur r_0 doit-on donner à r pour que $\varphi_1 = -\pi/2$?

2. Étude d'un filtre, d'après ENAC 2003

Soit le montage suivant où l'amplificateur opérationnel est idéal et fonctionne en régime sinusoïdal.

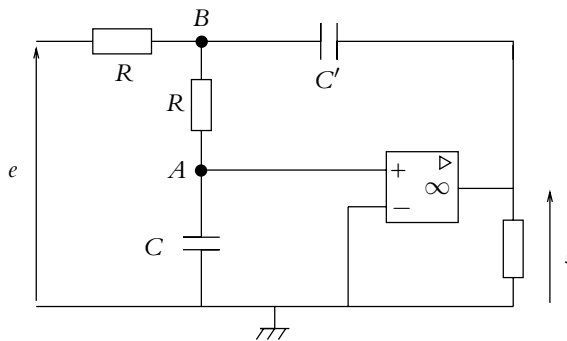


Figure 18.47

1. Établir l'expression de la fonction de transfert $\underline{H} = \frac{s}{\underline{e}}$.

1. $\underline{H} = \frac{1}{1 + R^2 C C' \omega^2 + j R C \omega}$

2. $\underline{H} = \frac{1}{1 + R^2 C^2 \omega^2 - j R C' \omega}$

3. $\underline{H} = \frac{1}{1 - R^2 C C' \omega^2 + 2 j R C \omega}$

4. $\underline{H} = \frac{1}{1 - R^2 C^2 \omega^2 - 2 j R C' \omega}$

2. Déterminer la relation entre les capacités C' et C pour que le gain s'écrive :

$$H = \frac{1}{\sqrt{1 + \frac{\omega^4}{\omega_0^4}}}$$

1. $2C' = C$

2. $C' = 2C$

3. $C = 4C'$

4. $C' = \sqrt{2}C$

3. En déduire l'expression de ω_0 .

1. $\frac{1}{4RC}$

2. $\frac{1}{RC\sqrt{2}}$

3. $\frac{2}{RC}$

4. $\frac{1}{2RC}$

4. Donner l'expression de l'asymptote à basses fréquences du diagramme de Bode en gain.

1. $G_{dB} = 0$

2. $G_{dB} = 20 \log_{10} \frac{\omega}{\omega_0}$

3. $G_{dB} = 40 \log_{10} \frac{\omega}{\omega_0}$

4. $G_{dB} = 40 \log_{10} \frac{\omega}{\omega_0}$

5. Même question à hautes fréquences.

1. $G_{dB} = 0$

2. $G_{dB} = 20 \log_{10} \frac{\omega_0}{\omega}$

3. $G_{dB} = 40 \log_{10} \frac{\omega_0}{\omega}$

4. $G_{dB} = 40 \log_{10} \frac{\omega}{\omega_0}$

6. Quel est le déphasage à basses fréquences ?

1. 0
2. π
3. $\frac{\pi}{2}$
4. $\frac{\pi}{4}$

7. Même question à hautes fréquences.

1. 0
2. π
3. $\frac{\pi}{2}$
4. $\frac{\pi}{4}$

8. Quelle est la nature du filtre considéré ?

1. actif
2. passe-bas
3. coupe-bande
4. passe-haut

3. Réjecteur de fréquence sélectif

Soit le montage :

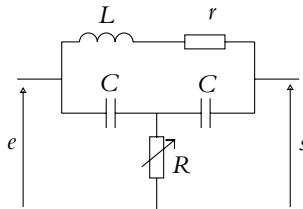


Figure 18.48

1. Déterminer la fonction de transfert à vide de ce filtre.

2. On souhaite obtenir un réjecteur de fréquence à la fréquence f_0 . Exprimer cette fréquence en fonction des valeurs des paramètres. En déduire la valeur R_0 de la résistance variable R permettant l'obtention d'un tel filtre.

3. On pose $x = \frac{f}{f_0}$ et $Q = \frac{2\pi f_0 L}{r}$. Écrire la fonction de transfert sous la forme $\frac{1}{1 + jA}$ en exprimant A en fonction de x et Q .

4. Tracer le diagramme de Bode du filtre.

5. Déterminer la bande passante à - 3 dB.

B. Exercices et problèmes

1. Diagramme de Bode d'un filtre du second ordre

1. Établir la fonction de transfert du filtre ci-contre sous la forme :

$$\underline{H}(j\omega) = \frac{\underline{f}(j\omega)}{\underline{f}(j\omega) + R_1 C_2 j\omega}$$

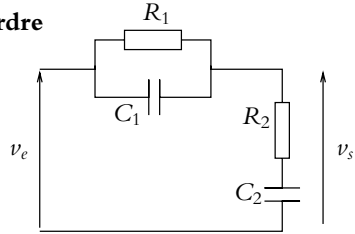


Figure 18.49

2. Montrer que l'on peut mettre le numérateur sous la forme :

$$(1 + j\tau_1\omega)(1 + j\tau_2\omega)$$

où $\tau_1 = R_1 C_1$ et $\tau_2 = R_2 C_2$ et le dénominateur sous la forme :

$$(1 + j\tau_3\omega)(1 + j\tau_4\omega)$$

où τ_3 et τ_4 sont à déterminer avec $\tau_3 < \tau_4$. On suppose :

$$(\tau_1 + \tau_2 + R_1 C_2)^2 > 4\tau_1\tau_2 \quad (18.26)$$

3. On pose $\lambda = \frac{C_1}{C_2}$, $\tau_1 = \tau = 10^{-3}$ s et $\tau_2 = 10\tau$.

Calculer λ pour avoir $\tau_3 = \frac{\tau}{5}$ et calculer τ_4 .

Vérifier avec les valeurs numériques que la condition (18.26) sur τ_1 et τ_2 est bien remplie.

4. Construire le diagramme de Bode en gain pour ce filtre par superposition des diagrammes des filtres du premier ordre mis en évidence précédemment.

2. Filtrés actifs

1. On étudie le montage ci-dessous avec un amplificateur opérationnel idéal. Déterminer le type de filtre et la fonction de transfert de ce filtre.

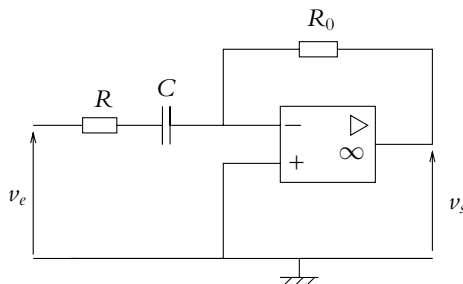


Figure 18.50

2. Calculer la fréquence de coupure ω_c avec $R_0 = 1$ M Ω , $R = 10$ k Ω et $C = 10$ nF.

3. On tient compte maintenant de la réponse en fréquence de l'amplificateur opérationnel avec gain en continu $\mu_0 = 2.10^5$ et un produit gain - bande passante de 3 MHz.

- Déterminer la pulsation de coupure ω_0 à vide de l'amplificateur opérationnel.
- Déterminer la nouvelle fonction de transfert du filtre en tenant compte des ordres de grandeurs de μ_0 , H_0 et $\frac{\omega_c}{\omega_0}$. De quel type de filtre s'agit-il ?
- Déterminer la bande passante de ce filtre.

3. Filtres à structure de Rauch, d'après École de l'air PSI 1999

Le problème étudie des montages ayant une structure de Rauch. Ces montages, selon le choix des impédances permettent d'obtenir pratiquement tous les filtres du second ordre. On note par $\underline{V}$ la représentation complexe de la grandeur V . On prendra $j^2 = -1$

1. Préliminaires

On réalise le montage ci-dessous où $\underline{Z}_i$ représente une impédance complexe.

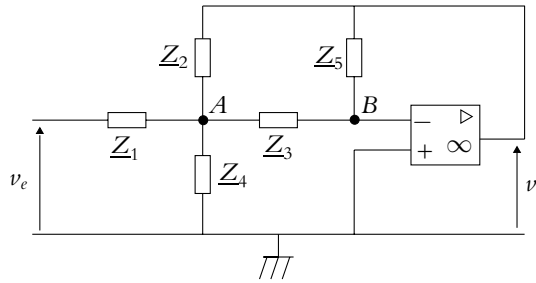


Figure 18.51

- Justifier rapidement que le montage est stable et que l'amplificateur opérationnel fonctionne en régime linéaire.
- Montrer que la fonction de transfert s'écrit :

$$\underline{H} = \frac{\underline{v}_s}{\underline{v}_e} = \frac{-\underline{Y}_1 \underline{Y}_3}{\underline{Y}_2 \underline{Y}_3 + \underline{Y}_5 (\underline{Y}_1 + \underline{Y}_2 + \underline{Y}_3 + \underline{Y}_4)}$$

$\underline{Y}_i$ représentant l'admittance complexe de l'impédance $\underline{Z}_i$.

2. Filtre passe-bande du second ordre

Les impédances $\underline{Z}_1$, $\underline{Z}_4$ et $\underline{Z}_5$ sont réalisées par des résistances ohmiques identiques R et les impédances $\underline{Z}_2$ et $\underline{Z}_3$ par des condensateurs identiques de capacité C .

- Déterminer la fonction de transfert. Montrer que cette fonction peut se mettre sous la forme :

$$\underline{H} = H_0 \frac{2j\lambda \frac{\omega}{\omega_0}}{1 + 2j\lambda \frac{\omega}{\omega_0} - \frac{\omega^2}{\omega_0^2}}$$

On précisera les valeurs de H_0 , λ et de ω_0 .

b) Représenter le circuit équivalent en basses fréquences (pour $\omega \ll \omega_0$) et établir la fonction de transfert $\underline{H}$ en basses fréquences.

Retrouver le résultat à partir de la fonction de transfert complète.

c) Représenter le circuit en hautes fréquences (pour $\omega \gg \omega_0$) et établir la fonction de transfert $\underline{H}$ en hautes fréquences.

Retrouver le résultat à partir de la fonction de transfert complète.

En déduire que il s'agit d'un filtre passe-bande.

d) Déterminer le module $|\underline{H}|$ de la fonction de transfert ainsi que sa phase. Le résultat sera exprimé en fonction des paramètres H_0 , λ et ω_0 .

e) Pour quelle pulsation ω_M le module $|\underline{H}|$ passe-t-il par un maximum H_M ?

On déterminera ω_M en fonction de ω_0 et on calculera la valeur numérique de H_M .

f) Déterminer les deux pulsations de coupure à -3 dB, ω_1 et ω_2 . Les deux pulsations seront exprimées en fonction de ω_0 . En déduire le facteur de qualité de ce circuit. On rappelle que

$$Q = \frac{\omega_M}{\omega_2 - \omega_1}.$$

g) Tracer le diagramme de Bode en gain et en phase. On précisera la pente des asymptotes et les coordonnées de leurs points d'intersection.

h) Dans quelle(s) domaine(s) de fréquence ce circuit peut-il représenter un circuit dérivateur?

i) Même question pour un circuit intégrateur.

4. Circuits à condensateur, d'après CCP TSI 2005

1. Charge d'un condensateur à travers une résistance :

On considère un circuit composé d'une résistance R et d'un condensateur de capacité C en série aux bornes duquel on place un générateur de tension idéal de force électromotrice constante E et un interrupteur K . Initialement le circuit est ouvert et le condensateur est déchargé. On note u_C la tension aux bornes du condensateur. A l'instant $t = 0$, on ferme l'interrupteur K .

a) Déterminer la valeur de l'intensité du courant parcourant le circuit au bout d'un temps infini.

b) Même question pour la tension u_C .

c) En déduire le comportement du condensateur dans ces conditions.

d) Établir l'équation différentielle vérifiée par u_C .

e) On pose $\tau = RC$. Quel nom donne-t-on à cette constante? Quelle est son unité dans le système international?

f) Établir l'expression de $u_C(t)$ et retrouver la valeur limite pour les temps longs.

g) Tracer l'allure de $u_C(t)$ en précisant les éventuelles asymptotes.

h) Donner l'équation de la tangente à l'origine.

- i) Calculer les coordonnées du point d'intersection de cette tangente avec l'asymptote.
- j) Déterminer, en fonction de τ , l'expression du temps t_1 à partir duquel la charge du condensateur diffère de moins de 1 % de sa charge finale.
- k) Donner l'expression de l'intensité parcourant le circuit en fonction du temps.

2. Étude énergétique de la charge du condensateur :

- a) Exprimer E_C l'énergie totale emmagasinée par le condensateur lors de sa charge en fonction de C et E .
- b) Même question pour E_J l'énergie dissipée par effet Joule dans la résistance lors de cette charge.
- c) Même question pour E_G l'énergie fournie par le générateur lors de cette charge.
- d) Donner une relation liant E_C , E_J et E_G et proposer une interprétation physique de cette relation.
- e) On définit le rendement énergétique de la charge du condensateur par $\rho = \frac{E_C}{E_G}$. Donner sa valeur ici.
- f) Pour améliorer ce rendement, on effectue la charge en deux phases : la première en chargeant le condensateur avec un générateur de force électromotrice $\frac{E}{2}$ puis en basculant sur un générateur de force électromotrice E une fois que C a atteint son maximum de charge sous une tension $\frac{E}{2}$.

Déterminer pour la première phase les expressions de E_{C1} et E_{G1} des énergies respectivement emmagasinée par C et fournie par le générateur.

- g) Établir l'expression de $u_C(t)$ lors de la seconde phase. On prendra pour nouvelle origine des temps l'instant de bascule entre les deux générateurs.
- h) En déduire celle de $i(t)$.
- i) Déterminer pour la seconde phase les expressions de E_{C2} et E_{G2} des énergies respectivement emmagasinée par C et fournie par le générateur.
- j) En déduire le nouveau rendement énergétique.
- k) Compte tenu des résultats obtenus, proposer une méthode pour faire tendre le rendement vers 1. On ne demande aucun calcul.

3. Circuit en régime sinusoïdal :

On reprend le circuit de la première partie en remplaçant le générateur de tension continue par un générateur basses fréquences délivrant une tension sinusoïdale d'amplitude E et de pulsation ω . On considère alors le filtre dont l'entrée est la tension du générateur basses fréquences et la sortie la tension aux bornes du condensateur.

- a) Comment se comporte le condensateur à hautes et basses fréquences ? En déduire la nature du filtre.
- b) Établir l'expression de la fonction de transfert du filtre et la mettre sous forme canonique.
- c) Étudier les variations du gain avec la pulsation.

- d) Même question pour le déphasage.
- e) Donner à basses fréquences les valeurs limites du gain, du gain en décibels et du déphasage.
- f) Même question à hautes fréquences.
- g) Donner les équations du diagramme asymptotique du gain en décibels.
- h) En déduire le tracé des diagrammes de Bode asymptotique et réel de ce filtre.
- i) Déterminer la fréquence de coupure de ce filtre.
- j) Quel type de comportement observe-t-on à hautes fréquences ?

4. Différents caractères d'un filtre :

Soit le montage suivant où l'amplificateur opérationnel est supposé idéal et fonctionne en régime linéaire.

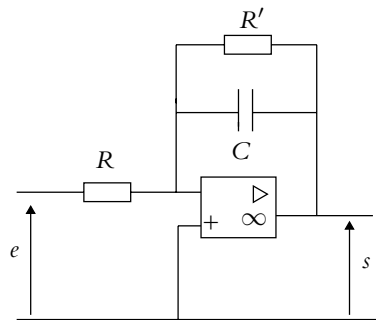


Figure 18.52

- a) Déterminer la fonction de transfert de ce montage.
- b) En déduire l'équation différentielle vérifiée par $s(t)$.
- c) À quelle condition sur R , C et ω ce montage a-t-il un comportement intégrateur ? Exprimer dans ce cas $s(t)$ en omettant les constantes d'intégration.
- d) À quelle condition sur R , C et ω ce montage admet-il une tension de sortie proportionnelle à celle d'entrée ? Justifier alors la dénomination d'amplificateur inverseur qui lui est donné. Préciser la valeur du gain ou coefficient d'amplification.
- e) Que peut-on dire de l'entrée inverseuse de l'amplificateur opérationnel ?
- f) Déterminer l'expression de $s(t)$ lorsque l'entrée est un échelon de tension entre 0 et E en supposant qu'initialement le condensateur est déchargé.
- g) En faisant un développement limité, établir qu'on peut observer un comportement intégrateur.
- h) Donner la condition sur R' , C et t pour que ce soit le cas.
- i) Quelle condition doit vérifier E pour que l'amplificateur opérationnel soit en régime linéaire ?

Étude expérimentale de quelques filtres

19

Ce chapitre est l'application expérimentale des résultats du chapitre précédent présentant la notion de filtrage et les outils permettant son étude comme la fonction de transfert et le diagramme de Bode. C'est également l'occasion de montrer l'importance du régime sinusoïdal forcé pour l'analyse de nombreux autres régimes. La décomposition des signaux en sommes de sinusoïdes par l'analyse de Fourier permet de ramener l'étude de tout système linéaire soumis à un signal périodique à celle du système soumis à un signal sinusoïdal. On peut alors déduire de l'étude du régime sinusoïdal l'influence d'un montage sur le contenu fréquentiel des signaux, ce qui traduit physiquement la notion de filtrage.

1. Décomposition harmonique d'un signal

1.1 Quelques observations

a) Spectre d'une sinusoïde

On a vu dans le chapitre 15 sur les circuits linéaires en régime sinusoïdal forcé qu'une sinusoïde $s(t) = A \sin(2\pi f_0 t)$ est caractérisée par sa fréquence f_0 et son amplitude A . On considère ici une phase initiale φ nulle ; si ce n'est pas le cas, il suffit d'opérer un changement d'origine des temps pour que cela le devienne. On peut donc représenter une sinusoïde de deux manières :

- représentation temporelle : il s'agit de l'allure du signal s en fonction du temps à partir de laquelle on peut déterminer sa fréquence f_0 et son amplitude A ;
- représentation fréquentielle : il s'agit de donner l'amplitude correspondant à la fréquence de la sinusoïde, c'est-à-dire tracer A en fonction de la fréquence.

Cette dernière représentation donne le spectre du signal, à savoir les fréquences qui le composent ainsi que leur importance relative par les amplitudes correspondant à chaque fréquence.

La figure suivante montre le cas où $f_0 = 50$ Hz et $A = 1$.

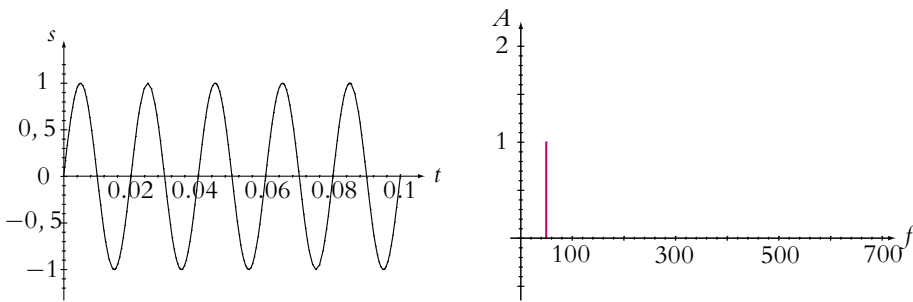


Figure 19.1 Signal sinusoïdal de fréquence 50 Hz et d'amplitude 1 et son spectre.

b) Cas de la somme de deux sinusoïdes

Si le signal $s(t)$ est la somme de deux sinusoïdes, on peut donner de la même manière le spectre de s :

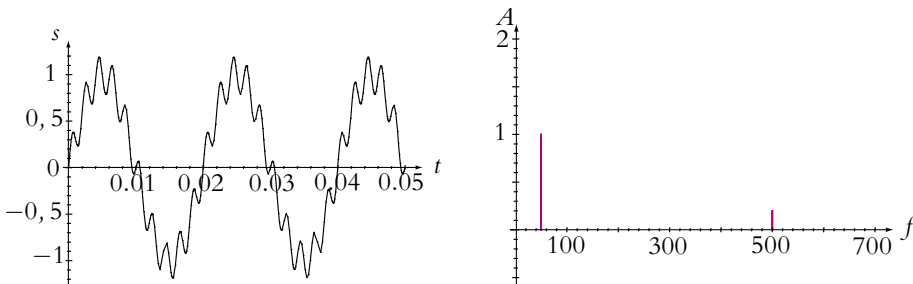


Figure 19.2 Allure du signal et spectre de la somme de deux sinusoïdes de fréquences respectives 50 Hz et 500 Hz et d'amplitudes respectives 1 et 0,2.

Les figures précédentes montrent ce qui se passe pour un signal :

$$s(t) = A_1 \sin(2\pi f_1 t) + A_2 \sin(2\pi f_2 t)$$

avec :

$$f_1 = 50 \text{ Hz}, f_2 = 500 \text{ Hz}, A_1 = 1 \text{ et } A_2 = 0,2.$$

On constate sur cet exemple simple qu'il est parfois plus facile de caractériser un signal par sa représentation fréquentielle que par sa représentation temporelle. Sur le spectre, il est aisé ici de voir quelles sont les fréquences présentes ainsi que leur importance relative *via* les amplitudes correspondantes. Sur la représentation temporelle, cette analyse est plus délicate et devient impossible dès lors que les fréquences constituant le signal sont en plus grand nombre.

D'autre part, l'allure temporelle n'est plus celle d'une sinusoïde bien que le signal reste périodique. Cette évolution conduit à considérer des signaux qui sont des sommes de sinusoïdes et en particulier le cas où les fréquences sont des multiples d'une fréquence f_0 .

c) Sommes de sinusoïdes et signaux périodiques

Soit un ensemble de fonctions sinusoïdales dont les fréquences sont des multiples d'une fréquence donnée f_0 .

La somme de ces fonctions pondérées par un choix judicieux de coefficients (qui sera précisé dans la suite de ce chapitre) permet de synthétiser des signaux périodiques simples ne ressemblant pas à une sinusoïde. La fréquence de ces signaux est la fréquence f_0 dont sont multiples toutes les fréquences des sinusoïdes utilisées.

Par exemple, on peut obtenir une dent de scie comme le montrent les figures 19.3. Les figures C1 à C10 donnent les sinusoïdes pondérées et les figures Sj, pour j variant de 1 à 10, la somme des contributions sinusoïdales des figures C1 à Cj. On constate qu'une dizaine de sinusoïdes permet d'obtenir de façon assez bonne une dent de scie. Il s'agit d'un résultat général qu'on va expliciter grâce aux séries de Fourier.

1.2 Fonction périodique et série de Fourier

Soit $f(t)$ une fonction périodique de $\mathbb{R}$ dans $\mathbb{R}$, de période T , de classe C^1 par morceaux¹. En notant a_n et b_n les coefficients de Fourier², le théorème de Fourier (qui sera vu dans le cours de mathématiques) permet alors d'écrire cette fonction :

¹ On rappelle qu'une fonction est de classe C^1 si elle est dérivable et que sa dérivée est continue.

² Les coefficients a_n et b_n sont définis par les relations :

$$\begin{cases} a_n = \frac{2}{T} \int_{t_0}^{t_0+T} f(t) \cos(n\omega t) dt \\ b_n = \frac{2}{T} \int_{t_0}^{t_0+T} f(t) \sin(n\omega t) dt \end{cases}$$

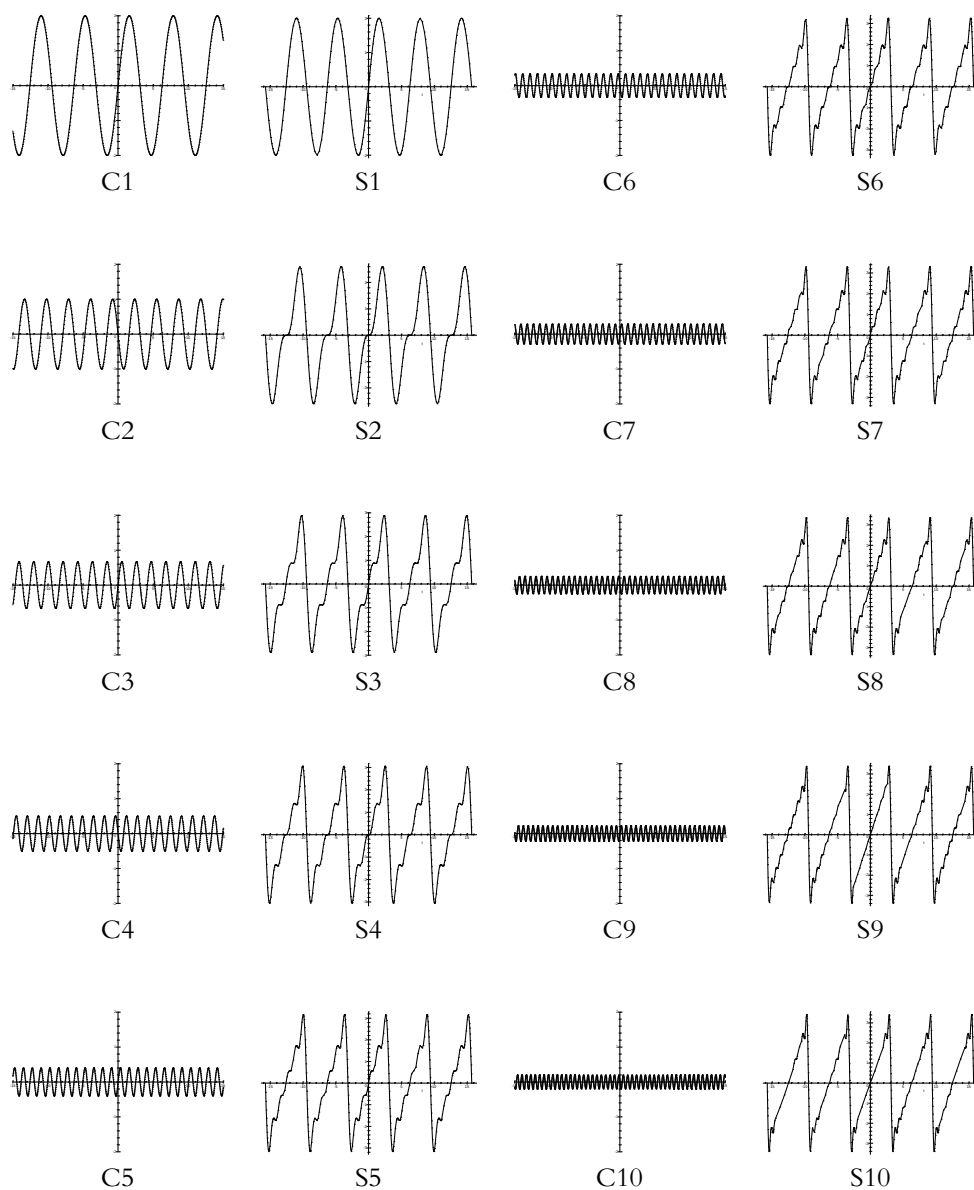


Figure 19.3 Synthèse d'une dent de scie par somme de sinusoïdes.

- en tout point où elle est continue comme une somme de termes sinusoïdaux :

$$f(t) = \frac{a_0}{2} + \sum_{n=1}^{+\infty} (a_n \cos(n\omega t) + b_n \sin(n\omega t)) \quad \text{avec} \quad \omega = \frac{2\pi}{T}$$

- en un point où elle est discontinue :

$$\frac{f(t^+) - f(t^-)}{2} = \frac{a_0}{2} + \sum_{n=1}^{+\infty} (a_n \cos(n\omega t) + b_n \sin(n\omega t))$$

Exemples de décompositions en séries de Fourier

a) Signal créneau

On adopte les notations suivantes pour le signal créneau :

$$f(t) = \begin{cases} A & \text{si } 0 < t < \frac{T}{2} \\ B & \text{si } \frac{T}{2} < t < T \end{cases}$$

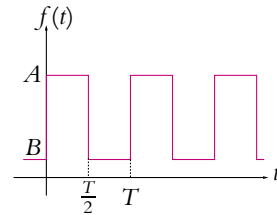


Figure 19.4 Signal créneau.

L'expression du développement en séries de Fourier de $f(t)$ est :

$$f(t) = \frac{A+B}{2} + \frac{2(A-B)}{\pi} \sum_{p=0}^{+\infty} \frac{\sin((2p+1)\omega t)}{2p+1}$$

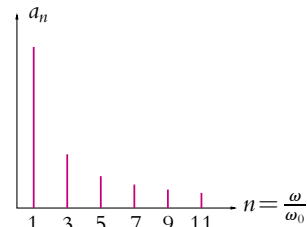


Figure 19.5 Spectre d'un créneau.

Si on ne prend que les premiers termes du développement, on obtient les sommes partielles qui ne donnent qu'une approximation de la fonction. Au fur et à mesure que le nombre de termes pris en compte augmente, la précision s'améliore. Par exemple, sur les figures suivantes, on a représenté les sommes partielles avec successivement 1, 3, 5, 7, 10, 25, 50 et 100 termes.

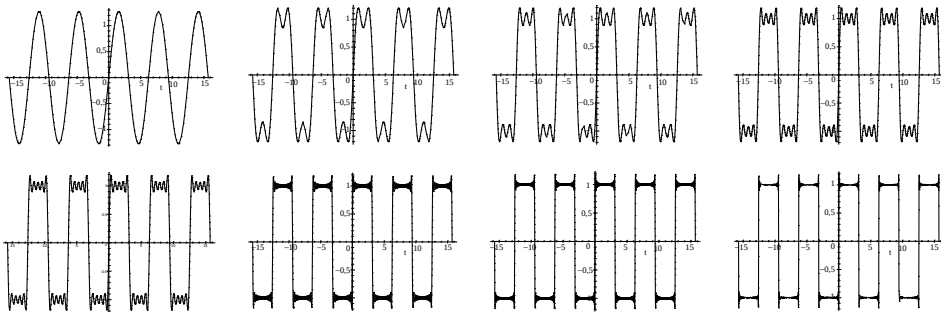


Figure 19.6 Sommes partielles d'ordre 1, 3, 5, 7, 10, 25, 50 et 100 d'un créneau.

Les figures ci-dessus correspondent au cas où $B = -A$.

b) Signal triangulaire

Soit un signal triangulaire de période T entre $B - A$ et $B + A$.

$$f(t) = \begin{cases} B + A \left(\frac{4}{T}t - 1 \right) & \text{si } 0 < t < \frac{T}{2} \\ B + A \left(-\frac{4}{T}t + 3 \right) & \text{si } \frac{T}{2} < t < T \end{cases}$$

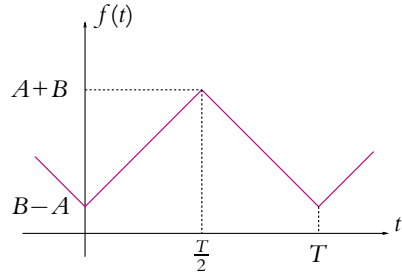


Figure 19.7 Signal triangulaire.

L'expression du développement en séries de Fourier de $f(t)$ est :

$$f(t) = B - \frac{8A}{\pi^2} \sum_{p=0}^{+\infty} \frac{\cos((2p+1)\omega t)}{(2p+1)^2}$$

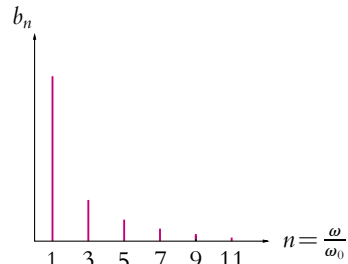


Figure 19.8 Spectre d'un signal triangulaire.

Les figures suivantes représentent les sommes partielles avec successivement 1, 3, 5, 7, 10, 25, 50 et 100 termes pour un signal triangulaire avec une valeur moyenne non nulle.

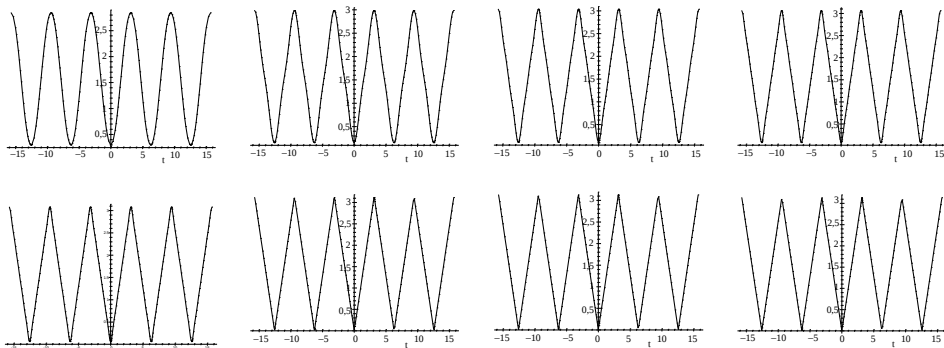


Figure 19.9 Sommes partielles d'ordre 1, 3, 5, 7, 10, 25, 50 et 100 d'un signal triangulaire.

On peut remarquer que le signal triangulaire se recompose plus facilement que le signal créneau. Ceci s'explique par la moins bonne régularité du créneau (il y a des points où le signal est discontinu) par rapport au triangle (les discontinuités concernent seulement la dérivée, le signal est continu). Il s'agit d'un résultat général : plus le signal présente de discontinuités, plus il faudra de termes dans les sommes partielles pour le reconstituer.

1.3 Analyse harmonique

À partir de la décomposition en séries de Fourier :

$$f(t) = \frac{a_0}{2} + \sum_{n=1}^{+\infty} (a_n \cos(n\omega t) + b_n \sin(n\omega t))$$

on peut donner les définitions suivantes.

a) Composante continue

Le coefficient a_0 correspond à la *composante continue* du signal, à savoir la partie du signal correspondant à la fréquence nulle. Il peut également être relié à la valeur moyenne du signal sur une période :

$$a_0 = \frac{2}{T} \int_{t_0}^{t_0+T} f(t) \cos(0 \times \omega t) dt = \frac{2}{T} \int_{t_0}^{t_0+T} f(t) dt = 2 \langle f(t) \rangle_T$$

Le coefficient a_0 est non nul si et seulement si le signal n'est pas centré autour de la valeur 0.

b) Fondamental et harmoniques

Le terme correspondant à $n = 1$ constitue le *fondamental* dont la fréquence est celle du signal. La pulsation ω est celle du fondamental et $f = \frac{\omega}{2\pi}$ est la fréquence du fondamental.

Pour $n > 1$, ce sont les *harmoniques d'ordre n* dont les fréquences sont des multiples de la fréquence du fondamental, à savoir de la fréquence du signal.

c) Analyse harmonique

Réaliser la décomposition en séries de Fourier d'un signal périodique consiste à rechercher les valeurs prises par les différents coefficients a_n et b_n pour la fonction $f(t)$ périodique considérée.

On dit qu'on effectue l'analyse harmonique ou l'analyse spectrale du signal lorsqu'on recherche les fréquences présentes dans la décomposition de Fourier et qu'on détermine les amplitudes correspondantes (données par les coefficients a_n et b_n). Il y a donc deux aspects à l'analyse spectrale :

- connaître les fréquences présentes dans le signal,
- déterminer leur importance relative *via* l'amplitude correspondante.

L'ensemble des coefficients a_n et b_n constituent le *spectre de Fourier* du signal.

On peut noter que la fréquence de base ou fréquence du fondamental est celle du signal périodique $f(t)$.

On peut également remarquer que le développement en séries de Fourier peut s'écrire sous la forme :

$$f(t) = \frac{a_0}{2} + \sum_{n=1}^{+\infty} A_n \cos(n\omega t + \varphi_n)$$

En effet, on a :

$$A_n \cos(n\omega t + \varphi_n) = A_n \cos(n\omega t) \cos \varphi_n - A_n \sin(n\omega t) \sin \varphi_n$$

soit en identifiant aux coefficients a_n et b_n :

$$\begin{cases} a_n = A_n \cos \varphi_n \\ b_n = -A_n \sin \varphi_n \end{cases}$$

1.4 Remarque sur le contenu spectral

Dans les exemples traités (signal carré ou signal triangulaire), on peut remarquer que les harmoniques d'ordre élevé ont des amplitudes plus importantes pour le signal carré que pour le signal triangulaire (la décroissance des coefficients est en $\frac{1}{n}$ au lieu d'être en $\frac{1}{n^2}$). Le créneau présente deux variations brutales au moment des discontinuités que les signaux sinusoïdaux ont du mal à rendre compte.

La décroissance des coefficients montre qu'il est possible, suivant le degré d'approximation qu'on accepte, de négliger les harmoniques d'ordre élevé sans que cela nuise fortement à la description du comportement. On parlera alors de *somme partielle d'ordre n* de Fourier lorsqu'on se limite aux harmoniques d'ordre inférieur ou égal à n .

Enfin, lorsqu'on compare les décroissances en $\frac{1}{n}$ pour le signal carré et en $\frac{1}{n^2}$ pour le signal triangulaire, on note qu'il faudra plus d'harmoniques pour le signal carré que pour le signal triangulaire pour un même degré d'approximation. Il s'agit d'une propriété générale : le contenu spectral est d'autant plus riche en harmoniques d'ordre élevé qu'on a des variations rapides. C'est logique sachant que les variations rapides correspondent à des fréquences élevées.

1.5 Généralisation

On s'est limité ici aux cas des fonctions périodiques. La notion d'analyse harmonique peut se généraliser aux fonctions non périodiques grâce à la transformée de Fourier. C'est ce qui sera vu dans le cours de seconde année.

2. Analyse harmonique et circuits linéaires

2.1 Importance de l'analyse harmonique pour les circuits linéaires

La propriété fondamentale des systèmes linéaires est le principe dit de superposition qui se traduit par le fait que la somme de solutions est également solution du problème. De ce fait, la décomposition de Fourier, qui permet d'écrire un signal sous la forme d'une somme de différentes composantes sinusoïdales, est particulièrement adaptée aux systèmes linéaires.

En effet, si on considère un signal d'entrée $e(t)$ périodique pouvant s'écrire sous la forme d'une série de Fourier, on pourra utiliser la procédure d'analyse suivante :

1. effectuer la décomposition en séries de Fourier de $e(t)$,
2. résoudre le problème pour chaque composante sinusoïdale de la décomposition en utilisant les complexes par exemple, et obtenir les différentes réponses de la décomposition en séries de Fourier,
3. déduire la solution en prenant la combinaison linéaire des différentes réponses.

On constate donc l'importance de la décomposition de Fourier pour les circuits linéaires. Cela justifie *a posteriori* l'étude des régimes sinusoïdaux du fait de leur rôle particulier dans les problèmes linéaires.

Cette idée est tout à fait générale : dans tous les domaines où on n'aura que des équations linéaires, la décomposition de Fourier et les réponses sinusoïdales auront ce rôle particulier. On pourra déterminer les réponses pour chaque pulsation et obtenir le résultat final par combinaison linéaire des différentes solutions trouvées. Cela justifie d'une part l'importance considérable des régimes sinusoïdaux en physique et d'autre part l'intérêt porté aux problèmes linéarisés pour lesquels cette technique peut être appliquée.

2.2 Lien avec les fonctions de transfert et les diagrammes de Bode

On rappelle que la fonction de transfert d'un système linéaire d'entrée $e(t)$ et de sortie $s(t)$ s'écrit à l'aide de la notation complexe :

$$\underline{H}(j\omega) = \frac{\underline{s}}{\underline{e}}$$

On s'était alors placé dans le cadre de signaux sinusoïdaux.

La décomposition en série de Fourier de tout signal périodique permet de généraliser l'utilisation de cette notion au cas de signaux périodiques mais non sinusoïdaux.

D'après la décomposition en séries de Fourier, tout signal d'entrée périodique $e(t)$ peut s'écrire sous la forme :

$$e(t) = \frac{a_0}{2} + \sum_{n=1}^{+\infty} (a_n \cos(n\omega t) + b_n \sin(n\omega t))$$

ou

$$e(t) = \frac{a_0}{2} + \sum_{n=1}^{+\infty} A_n \cos(n\omega t + \varphi_n) = \frac{a_0}{2} + \sum_{n=1}^{+\infty} e_{n\omega}(t)$$

Connaissant la fonction de transfert, on peut en déduire la réponse $s_{n\omega}(t)$ du système pour chaque composante sinusoïdale $e_{n\omega}(t)$, à savoir en notation complexe :

$$\underline{s}_{n\omega}(t) = \underline{H}(jn\omega) \underline{e}_{n\omega}(t)$$

et en revenant en notation réelle :

$$s_{n\omega}(t) = |\underline{H}(jn\omega)| A_n \cos(n\omega t + \varphi_n + \text{Arg}(\underline{H}(jn\omega)))$$

Dans le cas de systèmes linéaires, le principe de superposition permet alors d'écrire la réponse du système sous la forme de la combinaison linéaire suivante :

$$s(t) = \frac{s_0}{2} + \sum_{n=1}^{+\infty} s_{n\omega}(t) = \frac{s_0}{2} + \sum_{n=1}^{+\infty} A'_n \cos(n\omega t + \psi_n)$$

avec

$$\begin{cases} A'_n = |\underline{H}(jn\omega)| A_n \\ \psi_n = \varphi_n + \text{Arg}(\underline{H}(jn\omega)) \end{cases}$$

On obtient donc une réponse décomposée en séries de Fourier.

2.3 Utilisation du diagramme de Bode

Le diagramme de Bode qui donne une représentation graphique de l'influence du système sur l'amplitude et sur la phase de la sortie en fonction de celle d'entrée devient alors très pratique pour prédire les termes importants dans le développement du signal de sortie.

a) Analyse de l'amplitude

Pour déterminer l'influence du système sur l'amplitude, il faut connaître la position des fréquences du fondamental et de celles des harmoniques en fonction des fréquences caractéristiques apparaissant sur le diagramme de Bode en amplitude.

Lorsqu'une fréquence se trouve dans une zone où la réduction d'amplitude est forte, le terme correspondant dans le développement en séries de Fourier devient négligeable. Au contraire, si la réduction d'amplitude est nulle ou pratiquement nulle (le gain est alors proche de 1 et le gain en décibels proche de 0), la contribution reste importante.

Enfin si l'amplitude est augmentée, la contribution prendra plus d'importance. On ne conservera donc que les termes dont l'amplitude reste conséquente, ce seront souvent les termes pour lesquels le gain est proche de sa valeur maximale.

En ne conservant que les termes prépondérants, on obtient ainsi une sommation assez réduite pouvant permettre de prédire très rapidement l'allure qualitative du signal de sortie.

b) Cas de la composante continue

On a déjà dit que la composante continue correspond à une fréquence nulle. Le terme correspondant dans la décomposition en séries de Fourier est donc $\frac{a_0}{2}$. De ce fait, pour connaître l'influence du système sur la composante continue, il suffit de regarder sur le diagramme de Bode comment se comporte le gain en amplitude dans la limite des fréquences nulles et de raisonner comme pour les autres fréquences.

c) Analyse du déphasage

De la même manière, l'action du système sur la phase du signal pourra être déterminée en positionnant les fréquences du fondamental et des harmoniques sur le diagramme de Bode en phase.

Pour chaque fréquence sélectionnée lors de l'étude de l'amplitude, on déduit ainsi la valeur du déphasage introduit. On ne peut rien en déduire *a priori* sur l'allure du signal.

Si le déphasage est le même ou presque pour toutes les fréquences retenues au niveau de l'amplitude, il sera facile d'en déduire que le signal de sortie sera déphasé par rapport à celui d'entrée de cette valeur du déphasage. En revanche, si le déphasage varie énormément d'une fréquence à l'autre, il ne sera pas aisé de déterminer qualitativement l'allure du signal.

Finalement la première influence d'un système sur un signal est donnée par les modifications d'amplitude des différentes composantes spectrales. Il ne faut cependant pas négliger le déphasage introduit même s'il n'est pas forcément facile de prédire l'allure correspondante du signal.

3. Notion de filtrage

3.1 But du filtrage

Ce qui vient d'être vu conduit naturellement à la notion de filtrage. Un système linéaire invariant dans le temps conduit à une modification de la composante de pulsation ω :

- son amplitude est multipliée par un facteur égal au gain $G(\omega) = |H(j\omega)|$,
- sa phase augmentée d'un terme égale au déphasage $\varphi(\omega) = \text{Arg}(H(j\omega))$.

Comme G et φ sont des fonctions de la pulsation ω (ou de la fréquence $f = \frac{\omega}{2\pi}$), l'influence du système va évoluer avec la fréquence et par conséquent modifier le spectre de Fourier du signal de sortie par rapport à celui d'entrée. C'est sur cette évolution avec la fréquence qu'est basée la notion de filtrage.

On définit le filtrage comme la séparation des fréquences en fréquences dites utiles qu'on souhaite conserver dans le signal de sortie et en fréquences dites parasites qu'on désire voir disparaître. On doit alors utiliser un système linéaire invariant dans le temps tel que :

- le gain $G(\omega)$ soit significatif pour les fréquences utiles,
- le gain $G(\omega)$ soit négligeable pour les fréquences parasites.

La qualification utile ou parasite sera variable en fonction du but et de l'application poursuivie.

En termes de filtrage, on préfère les termes suivants :

- le système sera qualifié de filtre,
- la bande passante³ est définie comme la bande de fréquences transmises par le filtre ; on dira que le filtre laisse passer les fréquences correspondantes,
- la bande coupée correspond au complémentaire : il s'agit de l'ensemble des fréquences qu'on cherche à faire disparaître, on dit que la fréquence est coupée.

3.2 Filtre idéal

On dit qu'un filtre est idéal si le gain est constant dans la bande passante et vaut zéro dans la bande coupée.

Il est bien évident qu'en pratique le gain évoluera plus ou moins rapidement d'une valeur à l'autre.

3.3 Différents types de filtres

Suivant la zone de fréquences correspondant à la bande passante, on rappelle qu'on peut définir quatre types de filtres.

Leur analyse a été traitée au chapitre précédent ; le lecteur est invité à s'y reporter.

On rappelle que l'une des caractéristiques des filtres est la bande passante définie par

$$|\underline{H}(\omega)| \leq \frac{|\underline{H}|_{\max}}{\sqrt{2}} \quad \text{ou} \quad G_{dB}(\omega) \leq G_{dB,\max} - 3$$

La suite de ce paragraphe fait la synthèse des points essentiels de chaque type de filtres. Pour les figures, on a choisi $|\underline{H}|_{\max} = 1$.

³On verra qu'il s'agit de la même définition que celle donnée auparavant.

a) Filtres passe-bas

Ce type de filtre laisse passer les basses fréquences et coupe les hautes.

La bande passante est définie par $[0, f_c]$ en notant f_c la fréquence de coupure.

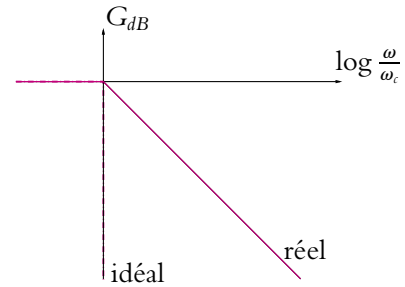


Figure 19.10 Filtre passe-bas.

b) Filtres passe-haut

Ce type de filtre est le complémentaire du précédent : il laisse passer les hautes fréquences et coupe les basses.

La bande passante est définie par $[f_c, +\infty[$ en notant f_c la fréquence de coupure.

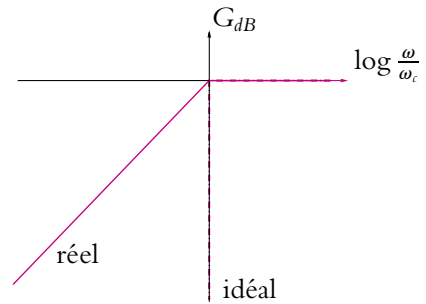


Figure 19.11 Filtre passe-haut.

c) Filtres passe-bande

Ce type de filtre coupe à la fois les hautes et les basses fréquences et ne laisse passer qu'une bande de fréquences $[f_1, f_2]$.

Cet intervalle définit la bande passante.

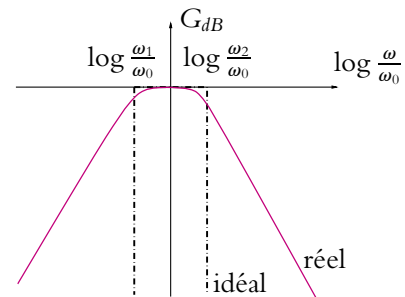


Figure 19.12 Filtre passe-bande.

d) Filtres coupe-bande ou réjecteur de fréquences

Ce type de filtre est le complémentaire des filtres passe-bande : il laisse passer à la fois les hautes et les basses fréquences et coupe une bande de fréquences $[f_1, f_2]$. Cet intervalle est la bande coupée.

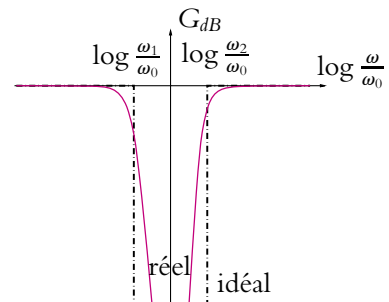


Figure 19.13 Filtre coupe-bande.

3.4 Ordre des filtres

Les filtres sont associés à une fonction de transfert $\underline{H}(j\omega)$ qui s'écrit comme une fraction rationnelle en $j\omega$ à savoir comme un quotient de deux polynômes en $j\omega$. On appelle *ordre du filtre* le degré du polynôme le plus élevé.

On remarque qu'il ne s'agit pas du degré de la fraction rationnelle $\frac{P}{Q}$ définie en mathématiques comme la différence du degré du numérateur et de celui du dénominateur $d = \deg P - \deg Q$.

Pour les filtres usuels, le degré du dénominateur est supérieur à celui du numérateur. L'ordre de ces filtres sera donc donné par le degré du polynôme en $j\omega$ du dénominateur.

On rappelle que l'importance des filtres du premier et du second ordre est due à la décomposition de tout filtre en filtres du premier et du second ordre : il est possible d'écrire toute fraction rationnelle sous la forme de produits de fractions rationnelles dont le numérateur et le dénominateur sont premiers entre eux et dont le degré est inférieur ou égal à 2.

4. Résultats du filtrage

4.1 Filtrage de la somme de deux sinusoïdes

On considère ici le signal

$$s(t) = \sin(2\pi f_1 t) + \sin(2\pi f_2 t)$$

avec $f_1 = 1$ kHz et $f_2 = 50$ kHz.

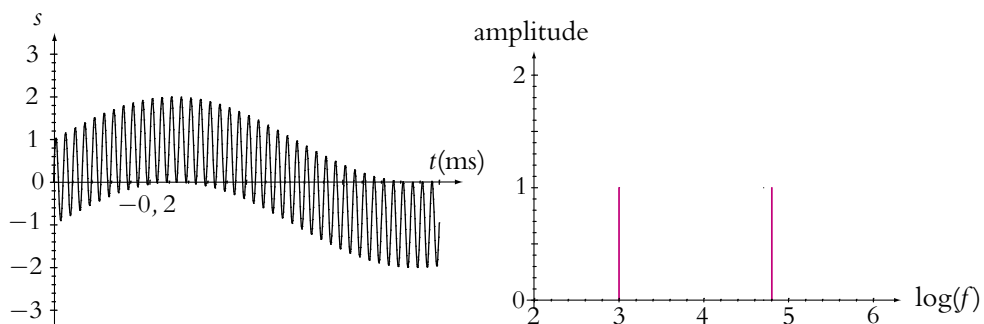


Figure 19.14 Évolution temporelle et spectre de la somme de deux sinusoïdes.

a) Filtrage par un filtre passe-bas

On envoie ce signal dans un filtre passe-bas du premier ordre de fréquence de coupure $f_c = 5$ kHz. En sortie du filtre, l'amplitude de f_1 devient 0,98 et celle de f_2 0,1. Le signal prend l'allure suivante :

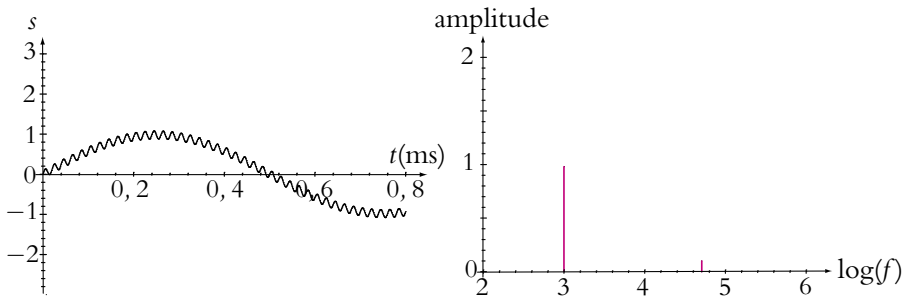


Figure 19.15 Évolution temporelle et spectre du filtrage de la somme de deux sinusoïdes par un filtre passe-bas dont la fréquence de coupure est entre les deux fréquences des sinusoïdes.

On constate donc que les hautes fréquences (ici f_2) sont quasi-inexistantes et que le signal est essentiellement composé de la partie à la fréquence f_1 comprise dans la bande passante.

b) Filtrage par un filtre passe-haut

On envoie ce signal dans un filtre passe-haut du premier ordre de fréquence de coupure $f_c = 10$ kHz. En sortie du filtre, l'amplitude de f_1 devient 0,1 et celle de f_2 reste à 1. Le signal prend l'allure suivante :

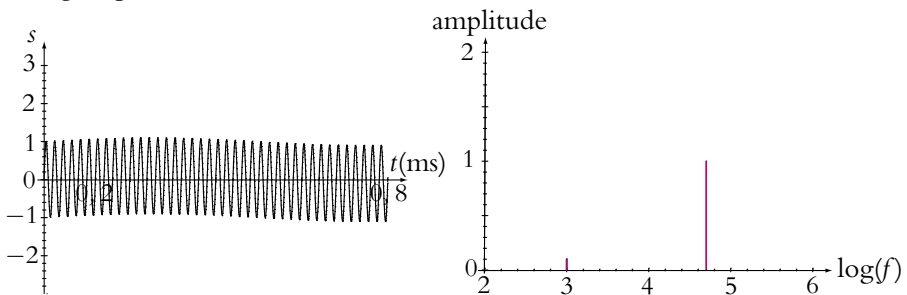


Figure 19.16 Évolution temporelle et spectre du filtrage de la somme de deux sinusoïdes par un filtre passe-haut dont la fréquence de coupure est entre les deux fréquences des sinusoïdes.

On constate donc que les basses fréquences (ici f_1) sont quasi-inexistantes et que le signal est essentiellement composé de la partie à la fréquence f_2 comprise dans la bande passante.

c) Filtrage par un filtre passe-bande (PCSI, PTSI)

On envoie ce signal dans un filtre passe-bande du second ordre de fréquence de résonance $f_0 = f_2 = 50$ kHz. Le facteur de qualité est pris égal à $Q = 20$ puis à $Q = 0,2$. Dans les deux cas, l'amplitude de la fréquence f_2 ou fréquence de résonance n'est pas modifiée. L'amplitude de f_1 , est négligeable (0,001) dans le premier cas et conserve une certaine importance (0,1) dans le deuxième. On constate ainsi l'influence du facteur de qualité sur le filtrage : un filtre passe-bande dont le facteur de qualité est bon sélectionne bien la fréquence souhaitée tandis qu'un filtre de facteur de qualité médiocre laisse passer une partie des composantes indésirables. Le signal prend l'allure suivante pour $Q = 20$:

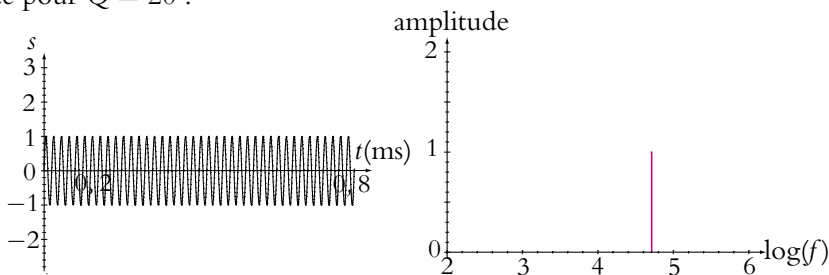


Figure 19.17 Évolution temporelle et spectre du filtrage de la somme de deux sinusoïdes par un filtre passe-bande de bon facteur de qualité et dont la fréquence de résonance est égale à la plus haute des fréquences des sinusoïdes.

et pour $Q = 0,2$:

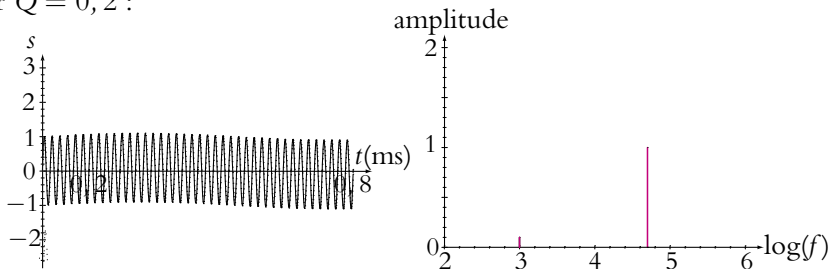


Figure 19.18 Évolution temporelle et spectre du filtrage de la somme de deux sinusoïdes par un filtre passe-bande de mauvais facteur de qualité et dont la fréquence de résonance est égale à la plus haute des fréquences des sinusoïdes.

On voit une légère ondulation de l'amplitude pour $Q = 0,2$ qui n'existe pas pour $Q = 20$.

d) Filtrage par un filtre coupe-bande (PCSI, PTSI)

On envoie ce signal dans un filtre coupe-bande de fréquence centrale $f_0 = f_2 = 50$ kHz. Le facteur de qualité est pris égal à $Q = 20$. L'amplitude de la fréquence f_2 est nulle et celle de f_1 vaut 1. Le signal prend l'allure suivante :

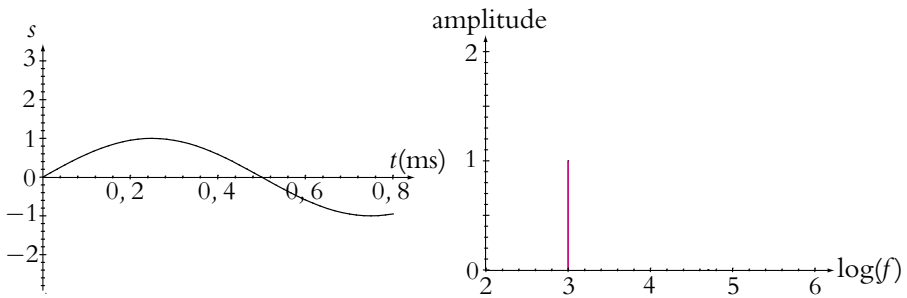


Figure 19.19 Évolution temporelle et spectre du filtrage de la somme de deux sinusoïdes par un filtre coupe-bande de bon facteur de qualité et dont la fréquence de résonance est égale à la plus haute des fréquences des sinusoïdes.

On n'a plus que la sinusoïde à la fréquence f_1 .

4.2 Filtrage d'un créneau

On fait la même analyse avec un créneau de fréquence f_0 représenté par les 20 premières harmoniques de son développement en séries de Fourier. Le signal d'entrée représenté sur la figure ci-dessous n'est pas un créneau idéal : il manque les harmoniques d'ordre élevé. On ne synthétise de même le signal de sortie qu'à partir des 20 premières harmoniques. De cette manière, on s'affranchit des écarts de reconstruction liés aux sommations partielles pour n'observer que l'influence du filtrage sur les signaux.

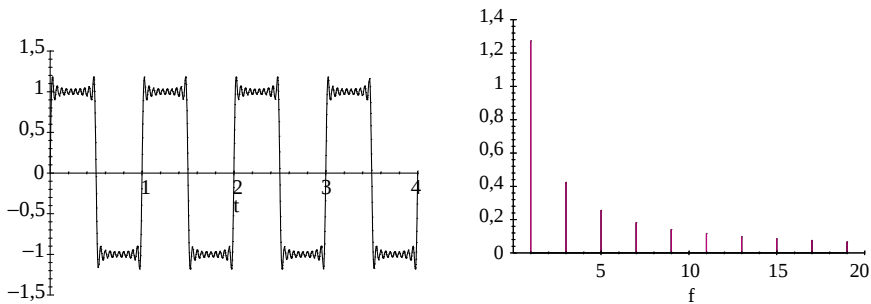


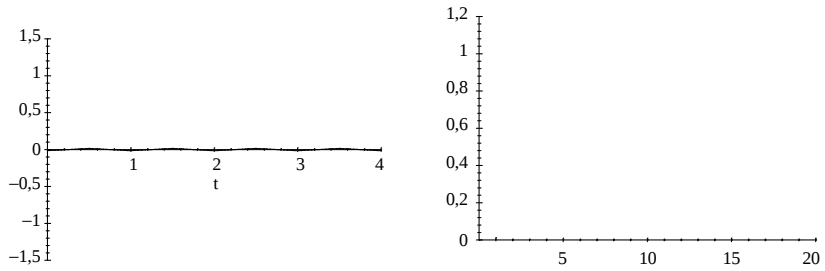
Figure 19.20 Créneau représenté par la somme partielle d'ordre 20 de son développement en séries de Fourier et son spectre.

On ne donne les résultats que pour un filtre passe-bas du premier ordre et un filtre passe-bande du second ordre. Le lecteur est invité à faire la même étude pour un filtre passe-haut du premier ordre et un filtre coupe-bande du second ordre en s'aidant des analyses réalisées pour la somme de deux sinusoïdes.

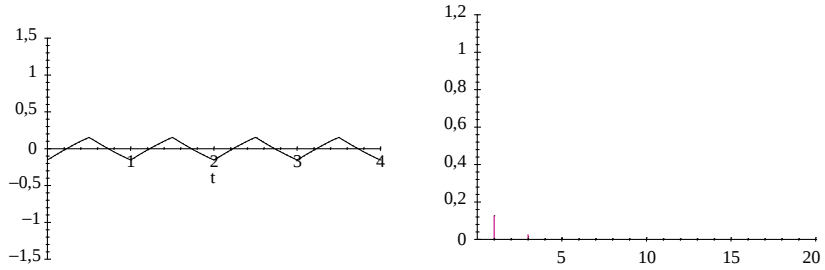
a) Filtrage par un filtre passe-bas

On utilise des filtres passe-bas du premier ordre dont les fréquences de coupure sont $0,01f_0$; $0,1f_0$; f_0 ; $10f_0$ et $100f_0$. Les allures obtenues sont les suivantes :

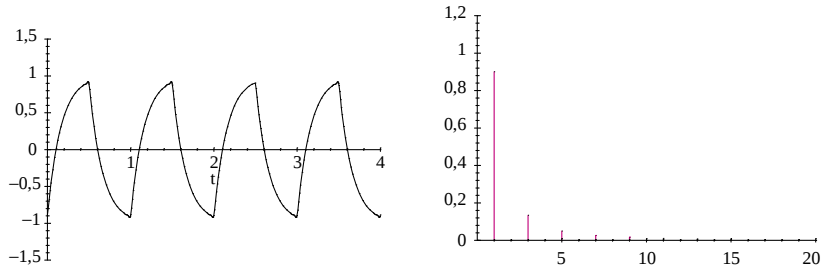
$$f_c = 0,01f_0$$



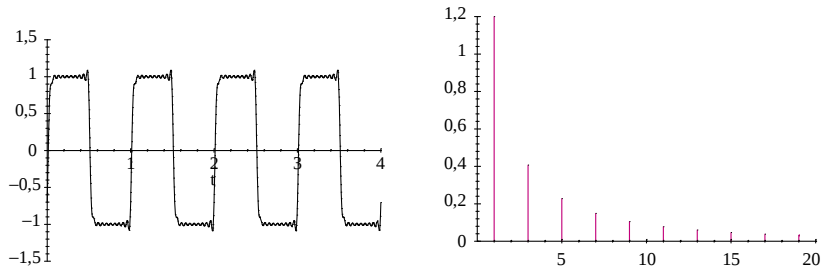
$$f_c = 0,1f_0$$



$$f_c = f_0$$



$$f_c = 10f_0$$



$$f_c = 100f_0$$

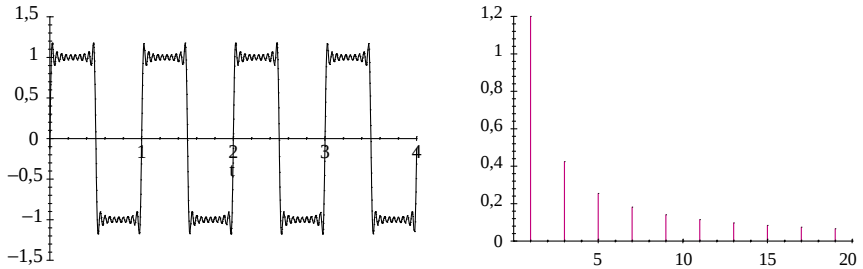


Figure 19.21 Filtrage d'un créneau de fréquence f_0 représenté par la somme partielle d'ordre 20 de son développement en séries de Fourier par un filtre passe-bas pour différentes fréquences de coupure f_c .

Lorsque la fréquence de coupure est inférieure à la fréquence f_0 du créneau (cas $f_c = 0, 01f_0$), le signal est inexistant ou presque.

Lorsque la fréquence de coupure se rapproche de f_0 (cas $f_c = 0, 1f_0$), on ne laisse passer que f_0 et ses tout premiers multiples : on obtient donc un signal triangulaire qui correspond à l'amorçage d'une sinusoïde. Le basculement du créneau entraîne une rupture qui se traduit par un point anguleux.

Quand $f_c = f_0$, on a le même phénomène avec une amplitude plus forte : le « poids » du fondamental et des harmoniques augmente.

Enfin, lorsque la fréquence de coupure est suffisamment grande ($f_c = 10f_0$ ou $f_c = 100f_0$), les fréquences que laisse passer le filtre permettent la reconstruction complète et sans déformation du créneau. Les amplitudes du fondamental et des harmoniques sont sensiblement les mêmes que celles du signal initial.

La sélection des composantes n'est pas très bonne du fait qu'on a un filtre du premier ordre : on a une pente à -20 dB par décade donc une diminution modérée des amplitudes. Avec un filtre d'ordre supérieur, ce serait meilleur.

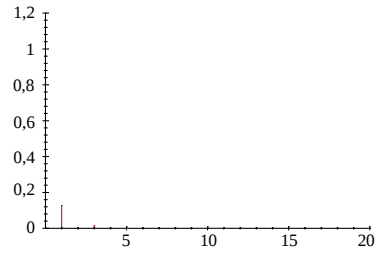
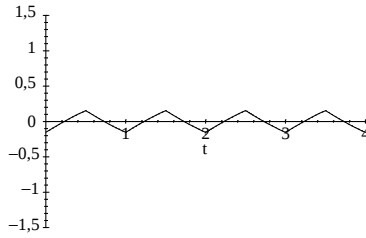
b) Filtrage par un filtre passe-bande (PCSI, PTSI)

On utilise des filtres passe-bande du second ordre dont les fréquences de résonance sont $0, 01f_0$; $0, 1f_0$; f_0 ; $10f_0$ et $100f_0$.

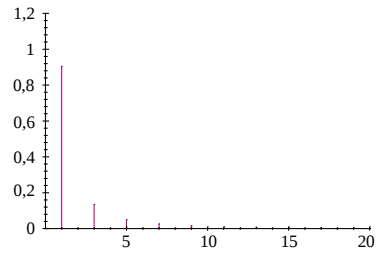
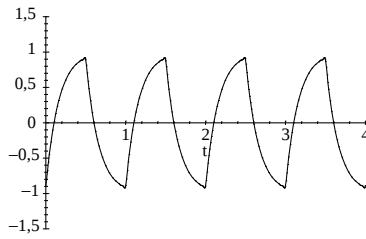
Suivant la valeur du facteur de qualité Q , on n'obtient pas les mêmes comportements.

Ainsi, pour un facteur de qualité $Q = 0, 1$, les allures obtenues sont les suivantes :

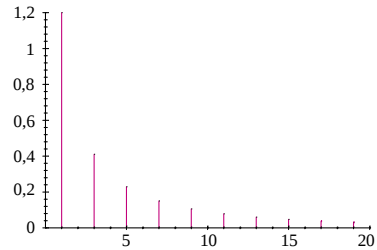
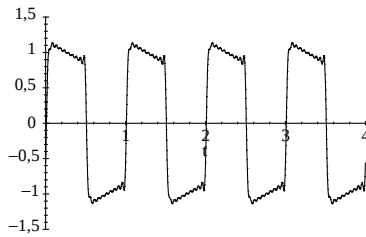
$$f_c = 0,01f_0$$



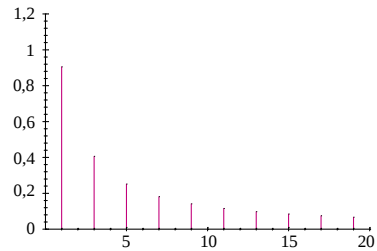
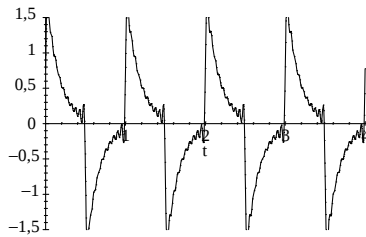
$$f_c = 0,1f_0$$



$$f_c = f_0$$



$$f_c = 10f_0$$



$$f_c = 100f_0$$

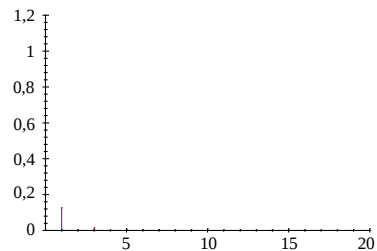
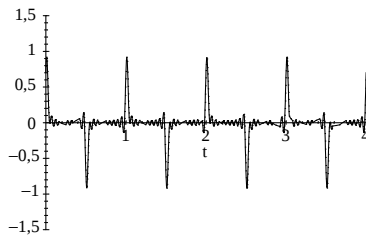


Figure 19.22 Filtrage d'un cr neau de fr quence f_0 repr sent  par la somme partielle d'ordre 20 de son d veloppement en s ries de Fourier par un filtre passe-bande de facteur de qualit  $Q = 0,1$ pour diff rentes fr quences de r sonance f_c .

On est dans le cas d'un facteur de qualité médiocre : il ne sélectionne pas très précisément une fréquence.

Ainsi, lorsque la fréquence de résonance est très inférieure à la fréquence f_0 du créneau ($f_c = 0,01f_0$), le signal est inexistant ou presque car toutes les composantes sont diminuées par rapport au signal initial.

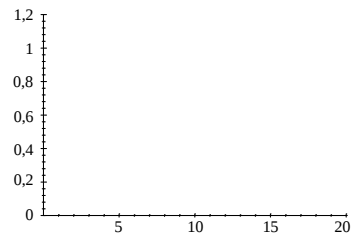
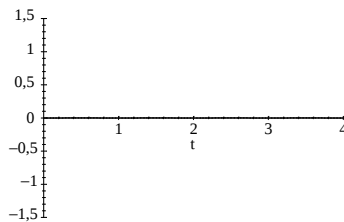
Lorsque la fréquence de résonance augmente et se rapproche de f_0 ($f_c = 0,1f_0$), on commence à laisser passer le fondamental et plus faiblement les harmoniques : on observe un triangle de faible amplitude comme avec un filtre passe-bas.

Quand la fréquence de résonance est proche de f_0 ($f_c = f_0$), le facteur de qualité est tel qu'il laisse passer le fondamental et quelques harmoniques : on obtient quasiment un créneau.

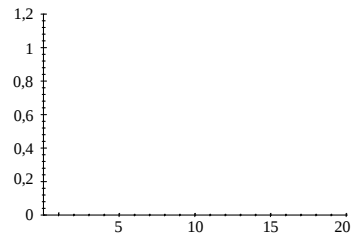
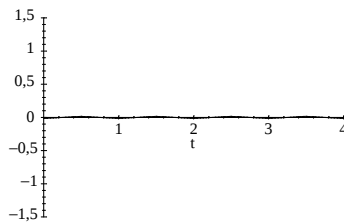
Puis quand on continue d'augmenter la fréquence de résonance ($f_c = 10f_0$ puis $f_c = 100f_0$), le fondamental sort de la bande passante puis les autres harmoniques également, on retrouve des allures analogues à celles d'un filtre passe-haut.

Si on considère maintenant un facteur de qualité meilleur $Q = 10$, on a les résultats suivants :

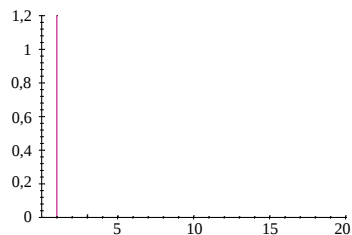
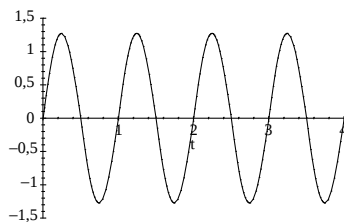
$$f_c = 0,01f_0$$



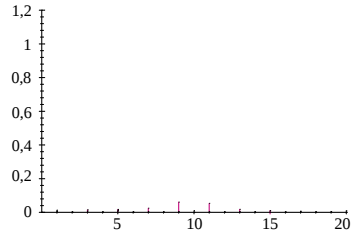
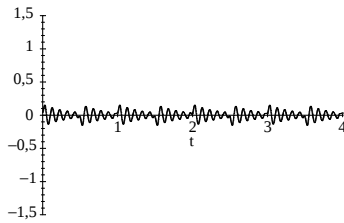
$$f_c = 0,1f_0$$



$$f_c = f_0$$



$$f_c = 10f_0$$



$$f_c = 100f_0$$

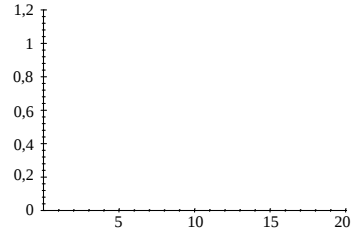
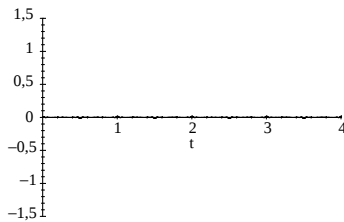


Figure 19.23 Filtrage d'un créneau de fréquence f_0 représenté par la somme partielle d'ordre 20 de son développement en séries de Fourier par un filtre passe-bande de facteur de qualité $Q = 10$ pour différentes fréquences de résonance f_c .

Le signal est inexistant lorsque la fréquence de résonance est différente de f_0 . Lorsqu'elle devient égale à f_0 , on n'a que le fondamental : les harmoniques sont en dehors de la bande passante et sont par conséquent coupées ; on observe alors une simple sinusoïde.

On notera que pour une fréquence de résonance valant dix fois la fréquence du créneau, on observe de faibles oscillations : des harmoniques ont une amplitude relativement significative pour donner lieu à ce comportement.

5. Étude d'un filtre sélectif

Dans ce paragraphe, on va donc s'intéresser à un filtre sélectif et aux effets de ce filtre sur les signaux. On considère le filtre suivant :

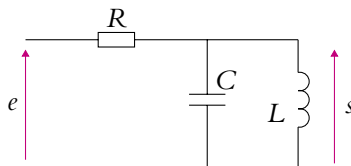


Figure 19.24 Montage du circuit bouchon.

On suppose qu'on peut négliger la résistance de la bobine : elle sera modélisée par une inductance pure.

5.1 Analyse théorique du comportement avec la fréquence

a) Étude des comportements limites

À très basse fréquence, la bobine se comporte comme un fil et le condensateur comme un interrupteur ouvert. La tension de sortie est donc nulle puisqu'elle est prise aux bornes de la bobine qui constitue un court-circuit. On a donc $H = 0$, le circuit ne laisse pas passer les basses fréquences.

À très haute fréquence, la bobine se comporte comme un interrupteur ouvert et le condensateur comme un fil. La tension de sortie est donc nulle puisqu'elle est prise aux bornes du condensateur qui constitue un court-circuit. On a donc $H = 0$, le circuit ne laisse pas passer les hautes fréquences.

L'impédance équivalente à la bobine (dont on néglige la résistance) et au condensateur associés en parallèle est :

$$\underline{Z}_{\parallel} = \frac{\frac{1}{jC\omega} jL\omega}{\frac{1}{jC\omega} + jL\omega} = \frac{\frac{L}{C}}{j\left(L\omega - \frac{1}{C\omega}\right)}$$

Lorsque $\omega = \omega_0 = \frac{1}{\sqrt{LC}}$, l'impédance devient infinie et on peut remplacer l'association bobine - condensateur par un interrupteur ouvert. La tension de sortie est alors celle de l'entrée : il n'y a pas de courant qui traverse la résistance R et donc pas de chute de potentiel entre l'entrée et la sortie. C'est également pour cette raison que ce montage est appelé *circuit bouchon* : l'intensité qui pourrait arriver par la résistance se trouve bloquée à la résonance par l'impédance infinie. Il se forme donc un « bouchon » de charges par analogie à ce que feraient des voitures devant une route fermée.

Le filtre laisse passer les fréquences au voisinage d'une pulsation $\omega_0 = \frac{1}{\sqrt{LC}}$ et coupe les hautes et les basses fréquences : il s'agit d'un filtre passe-bande.

b) Fonction de transfert

On suppose que la sortie de ce filtre est reliée à une charge d'impédance infinie : le courant de sortie est nul. On peut donc appliquer un pont diviseur de tension pour déterminer la fonction de transfert de ce filtre.

$$\underline{H} = \frac{\underline{s}}{\underline{e}} = \frac{\underline{Z}_{\parallel}}{R + \underline{Z}_{\parallel}} = \frac{1}{1 + j\frac{RC}{L}\left(L\omega - \frac{1}{C\omega}\right)}$$

On peut l'écrire sous la forme :

$$\underline{H} = \frac{1}{1 + jQ\left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega}\right)}$$

en notant $\omega_0 = \frac{1}{\sqrt{LC}}$ la pulsation propre et $Q = R\sqrt{\frac{C}{L}}$ le facteur de qualité.

On notera qu'on retrouve bien un comportement de type filtre passe-bande et que la pulsation propre est la même que celle du circuit R, L, C série. En revanche, le facteur de qualité est différent : il est égal à l'inverse de celui du circuit R, L, C série. De là provient l'intérêt de ce montage : si on souhaite augmenter le facteur de qualité Q , il faut par exemple augmenter la valeur de la résistance pour le montage traité ici. Dans le cas d'un circuit R, L, C série, il faudrait au contraire diminuer la valeur de la résistance avec le risque que cette valeur devienne alors proche de celle de la résistance interne du générateur de tension. Cela constitue donc une limitation importante à l'amélioration du facteur de qualité d'un circuit R, L, C série, ce qui n'est pas le cas ici.

L'allure du diagramme de Bode est le suivant (on se reportera aux chapitres antérieurs pour les justifications théoriques de cette allure) :

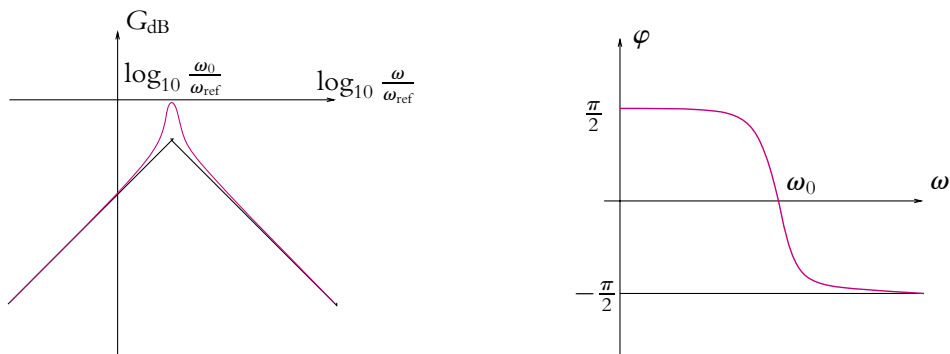


Figure 19.25 Diagramme de Bode du circuit bouchon.

5.2 Filtrage d'un créneau

On envoie en entrée un créneau dont on fait varier la fréquence f . D'après la décomposition en série de Fourier, on ne doit observer que les harmoniques $2\pi n f$ comprises dans la bande passante du filtre précédent : $\left[\omega_0 - \frac{\Delta\omega}{2}, \omega_0 + \frac{\Delta\omega}{2} \right]$ avec $\Delta\omega = \frac{\omega_0}{Q}$.

Pour une très faible valeur de f (petite devant $\frac{f_0}{2\pi} \left(1 - \frac{1}{2Q} \right)$ où $f_0 = \frac{\omega_0}{2\pi}$), les harmoniques qui ne sont pas supprimées par le filtre sont des harmoniques d'ordre élevé (par exemple, si $f = \frac{f_0}{100}$, ce seront les harmoniques dont l'ordre est au voisinage de 100). Leur amplitude initiale est très faible et en pratique, on n'observera que le régime transitoire. La figure ci-dessous représente ce qui se passe dans le cas où on a un régime pseudopériodique.

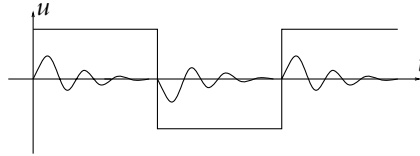


Figure 19.26 Filtrage sélectif d'un créneau de faible fréquence devant la fréquence de résonance du filtre.

Pour $f = f_0$, le fondamental appartient à la bande passante et, si on a un très bon facteur de qualité, les harmoniques d'ordre supérieur seront hors de la bande passante. On observe donc une sinusoïde de fréquence f_0 .

Au voisinage de la fréquence de résonance du filtre (entre $f_0 \left(1 - \frac{1}{2Q}\right)$ et $f_0 \left(1 + \frac{1}{2Q}\right)$), le fondamental appartient à la bande passante. On devrait donc également observer un signal sinusoïdal. Cependant les basculements d'une valeur à l'autre du créneau provoquent un régime transitoire qui se traduit par des discontinuités de la dérivée. On n'a donc que des bouts de sinusoïdes. La composante de fréquence $3f_0$ n'est pas négligeable en sortie du filtre, même si elle est plus faible que le fondamental (qui n'est pas amplifié au maximum car $f \neq f_0$).

À haute fréquence ($f \gg f_0 \left(1 + \frac{1}{2Q}\right)$), ni le fondamental ni les harmoniques n'apparaissent dans la bande passante. On n'aura donc que le régime transitoire qui donnera lieu à un signal triangulaire. En effet, à haute fréquence, le comportement du filtre est celui d'un intégrateur. L'amplitude sera cependant très faible.

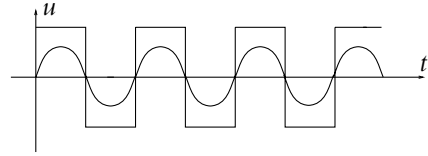


Figure 19.27 Filtrage sélectif d'un créneau dont la fréquence est la fréquence de résonance du filtre.

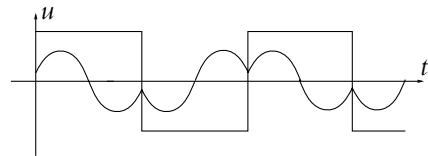


Figure 19.28 Filtrage sélectif d'un créneau dont la fréquence est proche de la fréquence de résonance du filtre.

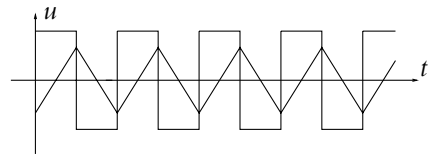


Figure 19.29 Filtrage sélectif d'un créneau de grande fréquence devant la fréquence de résonance du filtre.

A. Applications directes du cours

1. Action d'une filtre passe-haut sur un signal

On considère un filtre passe-haut du premier ordre dont la fréquence de coupure est 100 Hz. Donner l'allure du signal recueilli en sortie du filtre si on envoie en entrée :

1. une sinusoïde d'amplitude 4 V centrée autour de 1 V et de fréquence 2 kHz,
2. une sinusoïde d'amplitude 4 V centrée autour de 0 V et de fréquence 2 kHz,
3. un créneau d'amplitude 4 V centré autour de 1 V et de fréquence 2 kHz,
4. un créneau d'amplitude 4 V centré autour de 0 V et de fréquence 2 kHz.

2. Action d'une filtre passe-bas sur un signal

On considère un filtre passe-bas du premier ordre dont la fréquence de coupure est 100 Hz. Donner l'allure du signal recueilli en sortie du filtre si on envoie en entrée :

1. une sinusoïde d'amplitude 4 V centrée autour de 1 V et de fréquence 2 kHz,
2. une sinusoïde d'amplitude 4 V centrée autour de 0 V et de fréquence 2 kHz,
3. un créneau d'amplitude 4 V centré autour de 1 V et de fréquence 2 kHz,
4. un créneau d'amplitude 4 V centré autour de 0 V et de fréquence 2 kHz,
5. un créneau d'amplitude 4 V centré autour de 1 V et de fréquence 75 Hz,
6. un créneau d'amplitude 4 V centré autour de 0 V et de fréquence 75 Hz.

B. Exercices et problèmes

1. Filtre passe-bande, d'après CCP M 1994

On considère une bobine d'inductance L_1 et de résistance nulle qu'on associe en parallèle à un condensateur modélisé par une capacité C_1 de manière à constituer un circuit résonant. On place ce circuit $L_1 C_1$ dans le circuit collecteur d'un transistor à effet de champ convenablement polarisé de façon à réaliser un amplificateur sélectif. Le montage ainsi obtenu peut être représenté par le schéma équivalent suivant :

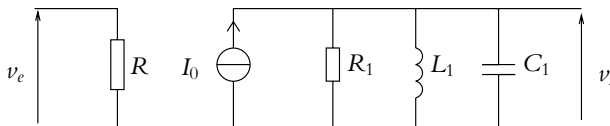


Figure 19.30

$\underline{I_0} = \frac{s \underline{v_e}}{R}$ et R_1 sont respectivement l'amplitude complexe du courant électromoteur et la résistance interne du générateur équivalent au circuit de sortie du transistor. s est supposé constant dans le domaine de fréquences considéré.

1. Exprimer la fonction de transfert $\underline{H} = \frac{v_s}{v_e}$ sous la forme : $\frac{A}{1 + jQ\left(x - \frac{1}{x}\right)}$ en notant $x = \frac{\omega}{\omega_0}$.
On exprimera A , Q et ω_0 en fonction de R_1 , L_1 , C_1 et s .
2. Quelle est la nature de ce filtre ?
3. Déterminer la bande passante de ce filtre.
4. On définit le taux d'harmonique de rang n d'un signal périodique par le rapport $\tau_n = \frac{V_n}{V_1}$ de l'amplitude V_n de l'harmonique de rang n à l'amplitude V_1 du fondamental. On note $\tau_{n,e}$ et $\tau_{n,s}$ les taux d'harmonique de rang n respectivement des signaux d'entrée et de sortie.
Calculer le taux d'atténuation $\delta_n = \frac{\tau_{n,s}}{\tau_{n,e}}$ de l'harmonique de rang n en fonction de Q et de n .
5. Pour quelle valeur de Q a-t-on un taux d'atténuation de -20 dB pour l'harmonique de rang 2 ? Même question pour un taux d'atténuation de -40 dB. On choisit pour la suite cette deuxième valeur du facteur de qualité.
6. La décomposition en séries de Fourier d'un créneau entre E_1 et E_2 s'écrit :

$$\frac{E_1 + E_2}{2} + \sum_{p=0}^{+\infty} \frac{2(E_1 - E_2)}{(2p+1)\pi} \sin(2p+1) \frac{2\pi}{T} t$$

Le signal d'entrée est un créneau de pulsation ω_0 . Quelle est l'allure du signal de sortie ? Justifier la réponse.

2. Filtre passe-bas à structure de Rauch, d'après École de l'Air PSI 1999

Ce problème est la suite de l'exercice B.3 du chapitre précédent. On s'y reportera pour les notations et les calculs préliminaires. Le lecteur est également invité à étudier le chapitre traité dans la suite sur l'amplificateur opérationnel dont il est fait usage ici.

On suppose dans ce problème que $\underline{Z}_1$, $\underline{Z}_2$ et $\underline{Z}_4$ sont des résistances identiques R et que $\underline{Z}_3$ et $\underline{Z}_5$ sont des condensateurs de capacité respective C_2 et C_1 . Les notations sont les mêmes que celles utilisées dans l'exercice B.3 du chapitre précédent.

1. Préciser dans ce cas la fonction de transfert qu'on notera $\underline{G}$ en fonction de R et C et montrer qu'on peut l'écrire sous la forme :

$$\underline{G} = \frac{G_0}{1 + 2jm\frac{\omega}{\omega_0} - \frac{\omega^2}{\omega_0^2}}$$

On explicitera les expressions de m , G_0 et ω_0 .

2. En déduire l'équation différentielle vérifiée par $s(t)$.
3. Que décrit la solution générale de cette équation ? Préciser les différentes allures possibles de cette solution générale.
4. On applique une tension constante E_0 à partir de l'instant $t = 0$, les deux condensateurs étant initialement déchargés. Donner, en les justifiant, les valeurs de $s(t)$ quand $t \rightarrow 0^+$ (notée $s(0^+)$) et quand $t \rightarrow \infty$ (notée $s(\infty)$). A quoi correspond physiquement la valeur de $s(\infty)$?

5. On veut obtenir $m = \frac{\sqrt{2}}{2}$. Calculer, en fonction de C_1 , la valeur qu'il faut donner à C_2 pour y parvenir. Quelle est alors la nature du régime libre ?
6. On veut faire varier la fréquence $f_0 = \frac{\omega_0}{2\pi}$ entre 1 kHz et 4 kHz. Entre quelles limites (exprimées en fonction de C_1) doit-on choisir R ? Dans la suite, on supposera que $f_0 = 2$ kHz.
7. Représenter en le justifiant l'allure du diagramme de Bode. On précisera les asymptotes et leur intersection.
8. Ce circuit peut-il être utilisé en intégrateur ? en dérivateur ?
9. On envoie un signal de fréquence $f = 1,5$ kHz.
 - a) Le signal est sinusoïdal centré. Représenter, en le justifiant, les signaux d'entrée et de sortie de ce filtre.
 - b) Mêmes questions pour un signal créneau pair de valeur basse 0 et de valeur haute a . On rappelle que la décomposition en séries de Fourier d'un tel signal est :

$$\frac{a}{2} + \frac{a}{\pi} \sum_{p=0}^{+\infty} \frac{(-1)^p}{2p+1} \cos((2p+1)2\pi ft)$$

On calculera éventuellement des ordres de grandeur pour justifier l'allure obtenue.

- c) Mêmes questions pour un signal triangulaire pair, centré, d'amplitude a .
10. Quel est l'avantage de ce filtre par rapport à un filtre passe-bas du premier ordre ?

3. Filtrage d'un signal électrique, d'après Banque PT 1999

Le signal électrique e issu d'un microphone est envoyé dans le circuit suivant :

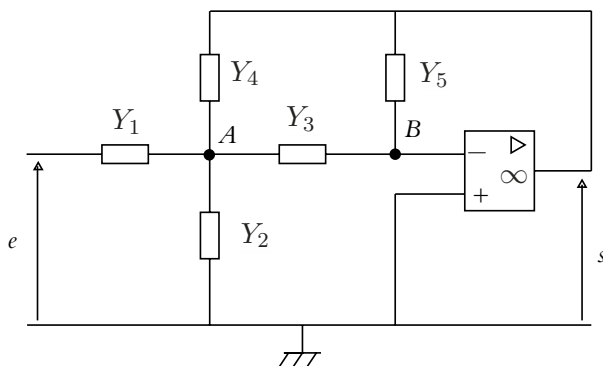


Figure 19.31

Les composants d'admittance Y_i sont soit des résistors soit des condensateurs.

1. Fonction de transfert générale :

Établir que la fonction de transfert de ce filtre peut se mettre sous la forme

$$\underline{H} = - \frac{\underline{Y}_1 \underline{Y}_3}{\underline{Y}_3 \underline{Y}_4 + \underline{Y}_5 \underline{S}}$$

On exprimera $\underline{S}$ en fonction des admittances $\underline{Y}_i$.

2. Etude du cas où $\underline{Y}_1 = \underline{Y}_3 = \underline{Y}_4 = \frac{1}{R}$, $\underline{Y}_2 = jC_2\omega$ et $\underline{Y}_5 = jC_1\omega$:

- Expliciter la fonction de transfert dans ce cas en fonction de R , C_1 , C_2 et ω .
- Quelle est la nature de ce filtre ?
- Écrire la fonction de transfert sous forme canonique. On précisera les valeurs des paramètres en fonction de R , C_1 et C_2 .
- Montrer alors qu'il est possible d'obtenir les valeurs des capacités connaissant la valeur de la résistance R et les caractéristiques du filtre.
- Si ce filtre est réalisé à l'aide de résistances dont les valeurs sont connues avec une précision infinie et qu'on a une précision de 5 % sur la détermination des caractéristiques du filtre, quelle est la précision (qu'on supposera identique) sur les mesures des capacités ?
- Application numérique : $R = 1,00 \text{ k}\Omega$, facteur de qualité $Q = 0,707$ et fréquence de résonance $f_0 = 20,0 \text{ kHz}$. Donner les valeurs numériques de C_1 et C_2 .

3. Réalisation pratique :

- On dispose d'un générateur basses fréquences de fréquence maximale 2 MHz et d'un oscilloscope deux voies ainsi que de tous les câbles nécessaires. Représenter le câblage à effectuer pour pouvoir tracer le diagramme de Bode du filtre.
- Décrire avec précision la manière dont on pourra réaliser ce tracé.
- On observe le chronogramme suivant pour un signal d'entrée particulier :

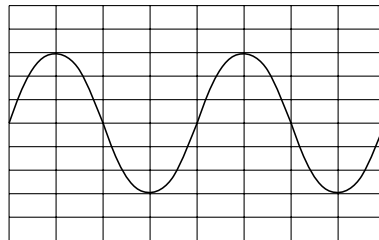


Figure 19.32

On a les calibres suivants : verticalement 1 V par division, horizontalement $12,5 \mu\text{s}$ par division. Quelle(s) est(sont) la(es) fréquence(s) de ce signal d'entrée et du signal de sortie correspondant ?

- d) Calculer avec les valeurs numériques précédemment fournies le module de la fonction de transfert ainsi que son déphasage.
- e) Représenter sur un même chronogramme les signaux d'entrée et de sortie.
- f) L'oscilloscope dispose d'un module permettant de tracer la décomposition en série de Fourier des signaux étudiés. Comment peut-on observer le diagramme de Bode en gain en utilisant un signal créneau ? On rappelle que l'analyse harmonique d'un créneau donne une décroissance en $\frac{1}{n}$ des amplitudes et qu'il est donc possible de considérer que les harmoniques de rangs élevés ont des amplitudes approximativement constantes.

4. Étude et utilisation de filtres, d'après ENSAIT 2002

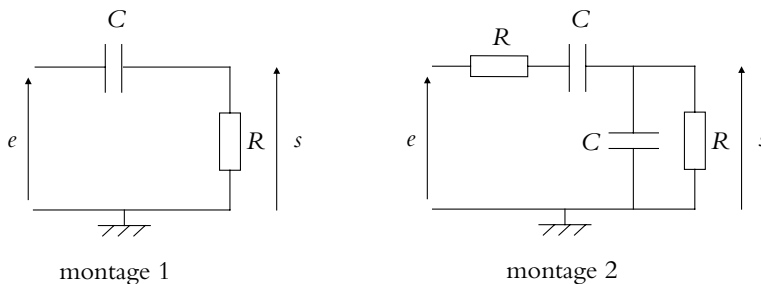


Figure 19.33

1. Déterminer la fonction de transfert $\underline{H}$ du montage 1.
2. Préciser la nature du filtre.
3. Définir le sens de variation du gain.
4. Donner l'expression du diagramme de Bode asymptotique en gain.
5. Déterminer la bande passante.
6. Établir l'expression du déphasage φ .
7. Déterminer son sens de variation.
8. Étudier les limites du déphasage φ .
9. Tracer le diagramme de Bode (asymptotique et réel).
10. Déterminer la valeur à donner à la capacité C pour que la tension de sortie soit affaiblie de 6 dB par rapport à la tension d'entrée lorsque celle-ci est un signal sinusoïdal de fréquence 1,0 kHz et que $R = 10,0 \text{ k}\Omega$.
11. Déterminer la fonction de transfert $\underline{H}'$ du montage 2.
12. Préciser la nature du filtre.
13. L'écrire sous la forme canonique.
14. Établir l'existence d'un phénomène de résonance dont on précisera la fréquence.

15. Établir l'expression du diagramme asymptotique en gain et préciser l'intersection des asymptotes.
16. Déterminer l'expression du déphasage ψ .
17. Établir le sens de variation du déphasage.
18. Donner les valeurs limites du déphasage.
19. Tracer le diagramme de Bode (asymptotique et réel).
20. En prenant les mêmes valeurs que précédemment pour R et C , donner la valeur de la fréquence de résonance.
21. Pour une fréquence $f = 1,0$ kHz, donner la valeur du gain pour le fondamental et les harmoniques d'ordre inférieure ou égal à 10.
22. Que peut-on en conclure sur l'influence du filtre sur un signal de fréquence f ?
23. Déterminer la bande passante et le facteur de qualité du filtre. Est-ce en accord avec le résultat précédent ?
24. On considère le filtre du montage 3. Sur quelle fréquence est-il accordé sachant que la fréquence d'accord est la fréquence de résonance en intensité du circuit R, L, C série ?

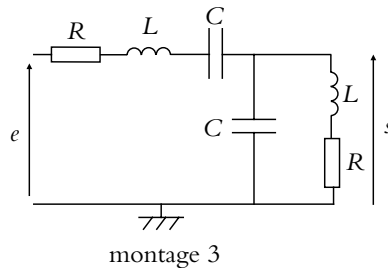


Figure 19.34

25. Calculer les impédances du circuit R, L, C série et du circuit R, L série en parallèle avec C à la fréquence d'accord.
26. On applique une tension sinusoïdale de valeur efficace 0,5 V et de fréquence la fréquence d'accord. Déterminer la tension de sortie du filtre à la fréquence d'accord.

20

T.P. Cours : Instrumentation électrique

On décrit dans ce chapitre l'utilisation des appareils électriques usuels aussi bien les générateurs que les appareils de mesure. On s'intéresse en outre aux caractéristiques des composants réels et aux écarts par rapport aux modèles théoriques qui ont été donnés précédemment. Enfin, on détaille quelques méthodes usuelles de mesure comme celles des déphasages et des impédances.

1. Générateurs basses fréquences ou GBF (PCSI, PTSI)

1.1 Principe

Il s'agit de dipôles actifs qui génèrent des signaux périodiques de différents types.

Ils sont essentiellement constitués d'un V.C.O. ou oscillateur de relaxation commandé en tension dont le principe sera évoqué dans le chapitre sur l'amplificateur opérationnel.

1.2 Types de signaux délivrés

La plupart des GBF délivrent des signaux sur deux sorties distinctes :

1. La sortie souvent notée « OUTPUT 50 Ω » sur laquelle il est possible en général de choisir l'un des signaux suivants :
 - signal sinusoïdal noté $\sim$,
 - signal triangulaire noté $\wedge$,
 - signal créneau noté $\sqcap$,
 - signal continu.
2. la sortie « OUTPUT TTL » ou « SYNC » qui délivre un signal créneau entre 0 et 5 V, les amplitudes ne pouvant pas être modifiées. Ce dernier type de signal est utilisé en logique T.T.L. (Transistor Transistor Logic) où on ne s'intéresse qu'à deux niveaux : un niveau bas ici, à 0 V, et un niveau haut, ici à 5 V. Cette

utilisation sera développée dans le cours de Sciences Industrielles. De tels signaux pourront être utiles en électricité pour la synchronisation des oscilloscopes qui sera détaillée plus loin.

1.3 Gammes de fréquences

On peut choisir la fréquence des signaux délivrés (sauf pour le signal continu pour lequel cette notion n'a pas de sens). On peut généralement faire varier la fréquence de 0 Hz à quelques MHz.

L'indication de fréquences donnée par l'appareil n'est en général pas précise (sauf si le GBF comporte un fréquencemètre et affiche la fréquence du signal délivré) : si on a besoin de la valeur de la fréquence, on fera une mesure par ailleurs.

1.4 Choix de l'amplitude

Le choix de l'amplitude correspond à la partie alternative et ne concerne donc pas les signaux continus. Ce choix ne concerne pas non plus les signaux T.T.L. par définition même de ces signaux. Les amplitudes peuvent varier de quelques millivolts à une dizaine de volts.

On dispose parfois aussi d'atténuateurs qui divisent l'amplitude des signaux par un facteur 10 ou 100.

Les GBF ne peuvent délivrer une tension supérieure en valeur absolue à une valeur limite $U_{\max}$. Elle est généralement de l'ordre d'une dizaine de volts. Cela signifie qu'on ne peut pas obtenir une tension en dehors de la bande $[-U_{\max}, U_{\max}]$. La conséquence de cette limitation est une éventuelle altération des signaux alternatifs. Par exemple, si on a un signal triangulaire de valeur moyenne 10 V et d'amplitude 6 V crête à crête, la tension à délivrer devrait varier entre $10 - 3 = 7$ V et $10 + 3 = 13$ V. Si le GBF ne peut délivrer plus de 12 V, il « écrête » le signal lorsque celui-ci dépasse 12 V. Il est donc possible d'obtenir des signaux écrêtés ayant, par exemple, l'allure suivante dans le cas de signaux triangulaires :

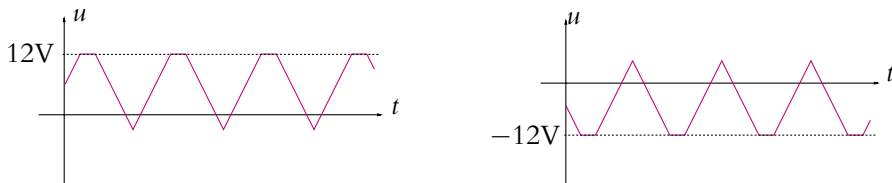


Figure 20.1 Saturation de la tension fournie par un GBF.

1.5 Tension de décalage ou d'offset

Une tension continue de décalage appelée aussi tension d'offset peut être ajoutée à la tension alternative pour les signaux de la sortie OUTPUT 50 Ω . La valeur de la tension d'offset est réglable dans la limite des tensions que le générateur peut délivrer.

1.6 Résistance interne

Les GBF ne sont pas des sources de tension idéales : ils possèdent une résistance interne de $50\ \Omega$, ce qui explique le nom donné à la sortie.

Il ne faudra jamais oublier l'existence de cette résistance, cela pourra notamment expliquer des différences entre le comportement prédit théoriquement et les observations lorsque la résistance interne du GBF ne pourra pas être négligée devant les autres impédances du circuit.

On peut mettre en évidence l'existence et l'importance de cette résistance interne en effectuant la manipulation suivante.

- On branche un oscilloscope ou un multimètre aux bornes du GBF : la tension lue est e_g , la tension aux bornes de la source idéale de tension car la résistance interne de $50\ \Omega$ du GBF est négligeable devant celle de l'appareil de mesure (de l'ordre au moins du $M\Omega$, Cf. paragraphes suivants). Il s'agit de la mesure à vide.
- On ajoute alors une résistance de charge R_c selon le branchement suivant :

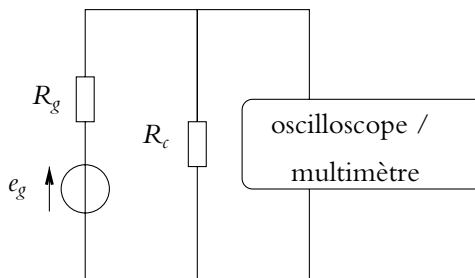


Figure 20.2 Mise en évidence de la résistance interne d'un GBF.

En supposant que R_c et R_g sont négligeables devant la résistance de l'appareil de mesure, l'intensité du courant qui traverse celui-ci est négligeable devant l'intensité du courant qui traverse R_c et R_g : ces deux résistances sont alors en série et on peut appliquer un diviseur de tension. On obtient :

$$e_{\text{mesuré}} = \frac{R_c}{R_c + R_g} e_g$$

Pour ne pas avoir à tenir compte de R_g , il faudrait que la tension mesurée soit égale à e_g . Ce sera le cas si $R_g \ll R_c$. Si la résistance de charge prend une valeur proche de la valeur de la résistance interne du GBF, il faudra tenir compte de R_g .

Cet exemple donne une méthode de mesure de R_g : on prend R_c variable et on la règle pour avoir $e = \frac{e_g}{2}$ (e_g ayant été mesurée « à vide » selon la procédure décrite plus haut). Dans ce cas, on a : $R_g = R_c$.

1.7 Balayage en fréquence ou wobulation

Les GBF comportent généralement une fonction de balayage en fréquence ou wobulation : cela permet de faire varier la fréquence du signal au cours du temps.

Si on choisit un signal sinusoïdal, le générateur fournit alors une tension :

$$u(t) = U_0 \cos(2\pi f(t)t)$$

où $f(t)$ varie entre deux valeurs extrêmes : $f_{\min}$ et $f_{\max}$.

Cette variation de $f(t)$ peut être linéaire :

$$f(t) = f_{\min} + \frac{f_{\max} - f_{\min}}{\Delta T} t$$

en notant ΔT la durée du balayage. Elle peut également être choisie logarithmique ou imposée par un signal extérieur.

Il est possible de régler la plage de fréquences balayées $f_{\max} - f_{\min}$ et la durée du balayage ΔT .

Le GBF fournit également une tension image de la fréquence, à savoir une tension fonction affine de la fréquence : il s'agit, dans le cas d'une wobulation linéaire d'une tension en dent de scie de durée ΔT , celle du balayage.

2. Alimentations stabilisées (PCSI, PTSI)

Une alimentation stabilisée est un dispositif fournissant une tension E continue indépendamment du courant ou une intensité I_0 indépendamment de la tension.

La caractéristique d'un tel générateur est la suivante :

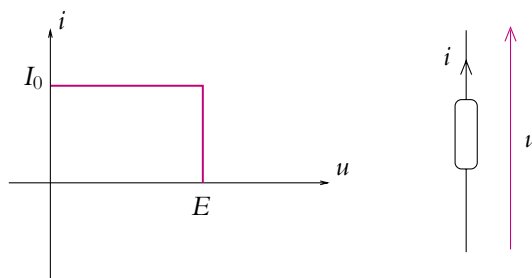


Figure 20.3 Caractéristique d'une alimentation stabilisée.

Les valeurs de E et de I_0 sont réglables dans une certaine plage, par exemple entre 0 et 30 V pour E et entre 0,1 et 2 A pour I_0 .

En pratique, la première chose à faire est de régler les valeurs E et I_0 .

On commence par I_0 en se plaçant en fonctionnement en générateur de courant. On relie alors les deux bornes de sortie : l'alimentation délivre un courant de court-circuit dont on peut régler la valeur. On choisit ainsi la valeur de I_0 .

On enlève le court-circuit et on bascule en fonctionnement générateur de tension. On peut alors régler la tension de l'alimentation stabilisée. On choisit ainsi la valeur de E .

Si on branche l'alimentation stabilisée aux bornes d'une résistance R , les valeurs de u et de i dans la résistance sont données par l'intersection de la droite $u = Ri$ et de la courbe $u = f(i)$ de la figure 20.3. On aura deux types de comportement :

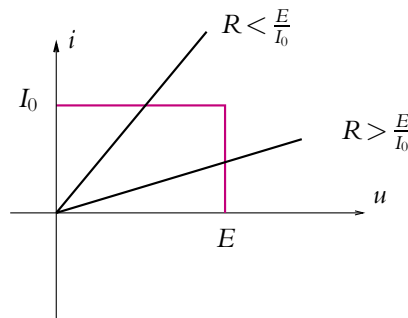


Figure 20.4 Les deux types de fonctionnement d'une alimentation stabilisée.

- Si $R < \frac{E}{I_0}$, l'alimentation impose un courant $i = I_0$ au montage et la tension vaut : $u = RI_0$.
- Si $R > \frac{E}{I_0}$, l'alimentation impose une tension $u = E$ au montage et l'intensité vaut : $i = \frac{E}{R}$.

3. Oscilloscopes

Il existe deux types d'oscilloscopes : les analogiques et les numériques.

3.1 Rappel : principe des oscilloscopes analogiques

Les oscilloscopes analogiques sont essentiellement constitués d'un tube cathodique et de circuits électroniques permettant de visualiser et d'analyser les signaux électriques.

Le tube cathodique (Cf. figure 20.5) comprend :

- un canon à électrons émettant un faisceau d'électrons quasi-homocinétiques : il fournira une tache lumineuse sur l'écran appelée « spot », la tension accélératrice de la cathode est classiquement de l'ordre de 2 kV,
- un système de focalisation permettant au spot d'être aussi fin que souhaité,
- des plaques de déviation entre lesquelles sont appliquées les tensions à analyser.

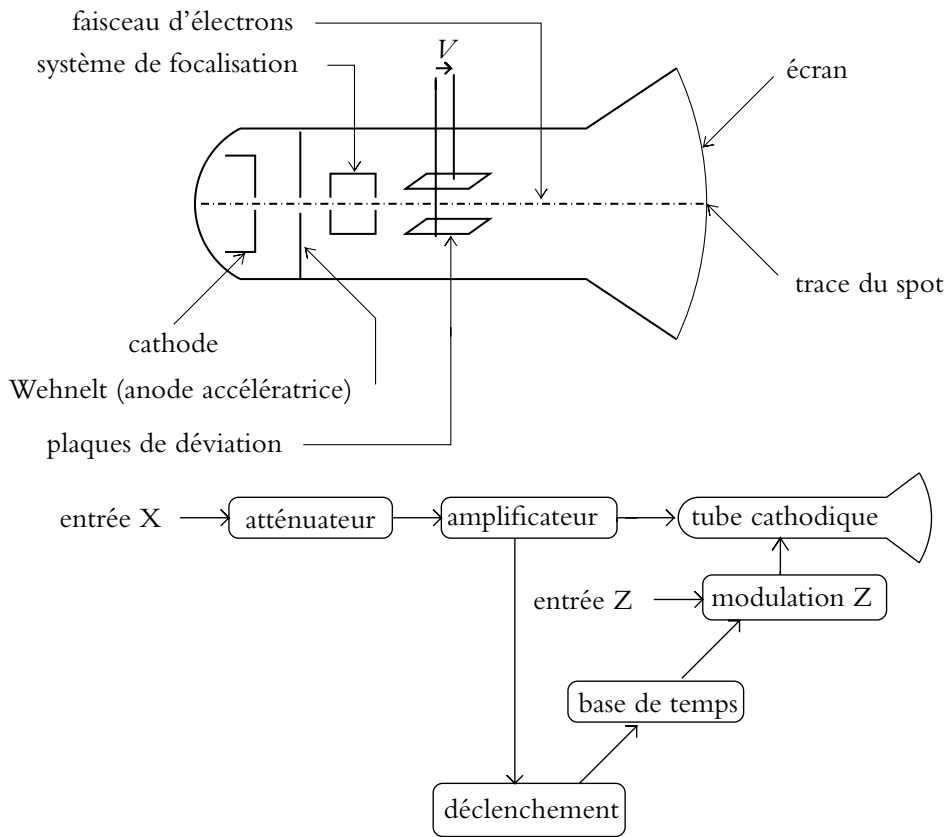


Figure 20.5 Schéma de principe de l'oscilloscope.

La tension V à analyser qui est appliquée entre les deux plaques de déviation distantes de d modifie la trajectoire des électrons.

La déviation s'exprime par :

$$x = -\frac{el^2}{2mdv_0^2} V$$

Le lecteur est invité à se rapporter au cours sur le mouvement des particules chargées pour la démonstration de cette expression.

La déviation est donc proportionnelle à la tension appliquée, ce qui permet de mesurer cette tension et de visualiser son évolution au cours du temps.

En dehors du tube cathodique, l'oscilloscope comprend différents circuits électroniques :

- un amplificateur réglable pour la tension étudiée V ,
- une base de temps,
- un circuit de déclenchement,

- un circuit Z de modulation de la luminosité du spot (par action sur le wehnelt sur lequel on reviendra plus loin).

Ces divers points seront analysés lors de la présentation des divers réglages.

3.2 Principe des oscilloscopes numériques

Une des étapes fondamentales du fonctionnement d'un oscilloscope numérique est la conversion analogique numérique. Elle comporte deux phases : l'échantillonnage et la quantification du signal.

Pour simplifier, échantillonner consiste à prendre la valeur du signal à des instants donnés (généralement à intervalle de temps régulier). On admet que le signal est correctement échantillonné si la fréquence d'échantillonnage est au moins deux fois plus grande que la fréquence maximale du signal (critère de Shannon). En pratique, on considère que 10 points par période suffisent à reproduire la forme du signal. Le risque en cas de non respect est, outre l'impossibilité d'obtenir la forme exacte du signal, de perdre des informations, comme des impulsions transitoires, ou d'observer des distorsions du signal en amplitude et en phase.

La base de temps numérique diffère de la base de temps analogique : elle correspond à l'intervalle de temps séparant deux points consécutifs. L'affichage étant généralement de 50 points par division, la base de temps en secondes par division correspond au produit de la base de temps numérique par 50. La base de temps numérique fixe la cadence de l'échantillonnage.

La quantification correspond à la traduction de la valeur mesurée à un instant dans un code binaire : le nombre de niveaux de valeurs possibles est pris égal à 2^N où N est le nombre de bits disponibles. En général, on a $N = 8$ soit 256 niveaux. La résolution est la plus petite variation de tension que l'appareil peut détecter : elle est donnée par le rapport entre la tension pleine échelle et le nombre de niveaux.

L'avantage d'utiliser un oscilloscope numérique vient du fait qu'il est possible d'ajouter des fonctions d'analyse et de traitement du signal aux fonctions habituelles des oscilloscopes. Parmi ces opérations supplémentaires, on peut citer la possibilité de faire des mesures, des moyennes en vue de réduire le bruit perturbant le signal utile, d'effectuer diverses opérations de filtrage numérique, d'effectuer différentes opérations mathématiques (transformée de Fourier discrète, intégration, dérivation...).

3.3 Réglages préliminaires d'un oscilloscope

Quand on allume un oscilloscope ou au cours d'une expérience, il peut être utile d'effectuer ou de modifier certains réglages comme le cadrage, la luminosité, la focalisation...

On peut déplacer la trace lumineuse verticalement ou horizontalement : cela correspond à des translations du faisceau d'électrons.

Le réglage de l'intensité lumineuse du spot consiste pour les oscilloscopes analogiques à choisir le débit des électrons à la sortie du canon à électrons en modifiant la tension du wehnelt. Cette tension varie classiquement entre 0 et -50 V .

Si le spot n'a pas une trace nette sur l'écran, il y a un problème de focalisation. Son réglage correspond pour les oscilloscopes analogiques à modifier la tension des anodes de focalisation qui varient classiquement entre 500 et 800 V. Cette opération doit être postérieure au réglage de la luminosité et sera donc effectuée chaque fois que l'intensité lumineuse sera modifiée.

La détermination de la sensibilité verticale correspond au choix du calibre ou de l'échelle des tensions étudiées. Elle est indiquée en mV ou V par division du cadran ; les calibres standardisés disponibles sont 1 mV ; 2 mV ; 5 mV ; 10 mV ; 20 mV ; 50 mV ; 0,1 V ; 0,2 V ; 0,5 V ; 1 V ; 2 V ; 5 V et parfois 10 V ; 20 V. On adapte le calibre aux variations des grandeurs étudiées. Il est également possible de « décaler » cette sensibilité.

3.4 Choix du circuit d'entrée

a) Modélisation du circuit d'entrée et influence sur les mesures

L'entrée de l'oscilloscope peut être modélisée électriquement, en couplage DC, par un dipôle passif constitué d'une capacité C_o et d'une résistance R_o en parallèle. Le rôle de la capacité est essentiellement un rôle de stabilisation de la tension mesurée.

Les paramètres d'entrée classiques d'un oscilloscope sont les suivants : $R_o = 1\text{ M}\Omega$ et C_o de 10 pF à 50 pF.

On va s'intéresser à l'influence de ces paramètres sur les mesures effectuées en régime continu : seule compte alors la résistance R_o . On effectue les opérations suivantes.

- On branche l'oscilloscope aux bornes d'un GBF : la tension lue est e_g , tension aux bornes de la source idéale de tension, car la résistance interne de $50\ \Omega$ du GBF est négligeable devant celle de l'oscilloscope.
- On ajoute alors une résistance de charge R_c selon le branchement suivant :

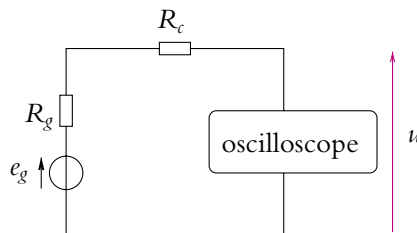


Figure 20.6 Mise en évidence de la résistance d'entrée de l'oscilloscope.

Dans ce cas, l'oscilloscope mesure :

$$u = \frac{R_o}{R_o + R_c + R_g} e_g$$

On pourra négliger l'influence de la résistance de l'oscilloscope si on retrouve la même tension que précédemment (en l'absence de résistance de charge) : il n'y a alors aucun courant significatif qui circule dans le circuit et la tension aux bornes des résistances R_c et R_g est négligeable. Ce sera effectivement le cas si $R_o \gg R_c$ et $R_o \gg R_g$. De façon générale, il faudra tenir compte de l'impédance d'entrée de l'oscilloscope si celle-ci a une valeur de l'ordre de grandeur des impédances du reste du circuit.

Pour mesurer les valeurs de R_o et C_o , on peut utiliser la procédure suivante :

- en continu ou en basse fréquence (100 Hz environ), on règle R_c pour avoir, par exemple, une tension mesurée de $\frac{e_g}{2}$ (e_g ayant été mesurée avec $R_c=0$) alors $R_o=R_c$,
- pour la détermination de C_o , on utilise un signal sinusoïdal : on choisit une fréquence du GBF telle que C_o n'ait pas une admittance négligeable devant $G_0 = \frac{1}{R_o}$ qu'on a déjà mesurée ; pour une fréquence donnée (par exemple 10 kHz) et une valeur donnée de R_c , on mesure alors le rapport des amplitudes de u et e_g pour en déduire C_o connaissant R_o .

b) Deux modes possibles : AC et DC

À l'entrée de l'oscilloscope, on dispose de trois modes : AC, DC et GND.

Le mode « GND (GROUND) » permet d'ajuster la référence du spot au potentiel 0 V. Cela revient à relier le signal du spot à la masse de l'oscilloscope.

Les deux autres positions doivent être clairement distinguées et on ne doit jamais choisir l'une ou l'autre au hasard.

La position « DC » correspond au cas où le signal étudié est envoyé sans intermédiaire à l'entrée de l'oscilloscope, ce qui explique le nom DC ou DIRECT CURRENT. On aura alors toutes les composantes du signal dont les fréquences sont comprises dans la bande passante de l'oscilloscope : on aura donc la composante continue ET la composante alternative. Par défaut, c'est cette position qu'il convient d'utiliser puisqu'on conserve la totalité de l'information sur le signal.

La position « AC » ou ALTERNATIVE CURRENT correspond à la suppression de la composante continue du signal, on ne garde que la partie alternative. Ceci est obtenu en plaçant un condensateur C' de grande capacité en série entre l'entrée de l'oscilloscope et le signal étudié. Ce composant est intégré à l'oscilloscope et sa mise en place est obtenue en choisissant le mode AC. Le circuit présente alors une fréquence de coupure de l'ordre de 2 Hz : les fréquences inférieures à 2 Hz sont éliminées.

On peut mettre en évidence ces différences par les manipulations suivantes.

- On observe à l'oscilloscope un signal sinusoïdal auquel on ajoute une tension continue. En mode DC, on observera une sinusoïde décalée de la tension continue alors qu'en mode AC, elle sera centrée autour de la valeur 0.

- On observe un signal crête à crête de fréquence relativement faible (50 Hz par exemple). En mode DC, on observe le crête à crête tel qu'on l'attend alors qu'en mode AC, le signal est déformé. On a des bouts d'exponentielles au lieu d'avoir des droites horizontales, ces exponentielles correspondent à des charges et des décharges du condensateur qu'on ajoute lors de l'utilisation du mode AC.

3.5 Modes de fonctionnement

Le choix du mode de fonctionnement correspond à la définition des voies qui vont être visualisées et de la façon dont l'oscilloscope va les représenter.

a) Fonctionnement unicourbe

Dans ce cas, on opte pour la position X ou Y correspondant à la voie sur laquelle seront effectuées les observations.

b) Fonctionnement bicourbe

Il est possible de visualiser simultanément à l'écran deux signaux. Il suffit, pour les oscilloscopes analogiques, de se placer sur l'une des positions ALT ou CHOP généralement. Pour CHOP, le spot trace un bout de la voie 1, puis un bout de la voie 2, puis un bout de la voie 1, puis un bout de la voie 2, etc. Pour ALT, le spot trace la voie 1 sur toute la largeur de l'écran, puis la voie 2 sur toute la largeur de l'écran.

Il est souvent possible de visualiser l'opposé d'une voie au lieu de la voie elle-même.

c) Mode XY

Lorsque le mode XY est sélectionné, on visualise la voie 2 en fonction de la voie 1. L'intérêt de ce mode sera montré lors des mesures de déphasages par exemple. Certains réglages comme les positions horizontale ou verticale (à savoir les réglages des zéros) sont en général à refaire lorsqu'on passe du mode bicourbe au mode XY ou l'inverse.

3.6 Synchronisation

a) Balayage

Pour observer une tension en fonction du temps, on applique sur les plaques de déviation horizontale une tension en dent de scie du type de celle représentée sur la figure 20.7. Entre t_i et $t_i + T$, on a une fonction linéaire du temps fournissant un balayage horizontal du spot linéaire. La trace parcourt l'écran de gauche à droite au cours d'un balayage.

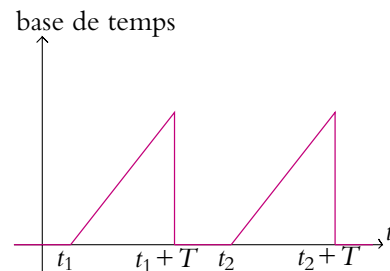


Figure 20.7 Allure de la base de temps.

b) Nécessité de synchroniser

Le signal visualisé à l'écran est la superposition des traces produites au cours des balayages successifs.

Le balayage en dents de scie implique que la trace se forme entre t_1 et $t_1 + T$ puis t_2 et $t_2 + T$ etc. Pour obtenir une SEULE trace sur l'écran, il faut que les instants $t_1, t_2, \dots$ correspondent rigoureusement au même point du signal. Ceci doit être réalisé en choisissant correctement la base de temps. On dit alors que la base de temps est *synchronisée* avec le signal à analyser.

Dans le cas contraire, on a soit une figure inextricable, soit une figure glissante.

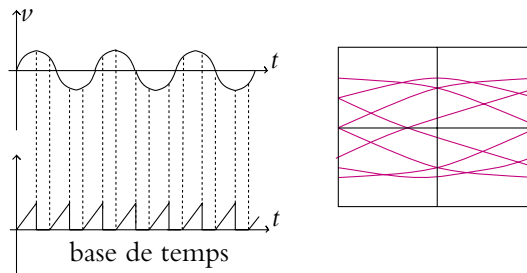


Figure 20.8 Illustration de la nécessité de synchroniser la base de temps et le signal.

On peut avoir deux types de balayage :

- balayage relaxé : le balayage est ininterrompu ($t_2 = t_1 + T$). Dans ce cas, l'oscilloscope ne sera synchronisé que si T est un multiple de la période du signal, ce qui est rare... Le seul réglage possible est celui de la fréquence de balayage.
- balayage déclenché : le spot se déplace sur l'écran pendant le temps T appelé durée du balayage et il est éteint entre deux balayages successifs. Cette extinction est assurée par le circuit Z. Il faut alors que les différentes traces débutent quand le signal est revenu dans une situation analogue à celle du premier déclenchement. Ceci s'obtient en choisissant la source de déclenchement et en réglant le niveau de ce dernier.

c) Choix de la fréquence de balayage

On peut régler la fréquence de la base de temps à l'aide de la molette « T/DIV ». Il faut noter que cela définit également l'échelle des abscisses à savoir l'échelle de temps. On a en général 18 positions accessibles entre $0,5 \mu\text{s}$ et 200 ms par division.

Comme pour le réglage de la sensibilité verticale, il est possible de décalibrer la base de temps. On verra l'utilisation de cette fonctionnalité pour les mesures de déphasage.

d) Choix de la source de déclenchement

La source de déclenchement, qui est le signal à partir duquel se déclenche le balayage, peut être couplée avec différents signaux. Il s'agit de comparer ce signal de déclen-

chement à une tension de référence interne, appelée niveau de déclenchement ou LEVEL, qui est constante et réglable. À l'instant où le signal de déclenchement atteint le niveau de déclenchement, le déclencheur envoie une impulsion à la base de temps pour que le balayage débute. Quand le balayage est terminé, le spot reste en attente jusqu'à ce qu'à nouveau le signal de déclenchement atteigne le niveau de déclenchement, ce qui le fait redémarrer.

Plusieurs choses doivent être définies :

1. la source de déclenchement : il est possible de déclencher sur l'une des deux voies X ou Y, sur un signal extérieur EXT envoyé sur la borne correspondante de l'oscilloscope, le réseau électrique (50Hz) LINE ou le signal propre à chaque voie ALT. Dans ce dernier cas, l'oscilloscope déclenchera sur la voie 1 pour les visualisations X, CHOP et ADD, sur la voie 2 pour les visualisations Y et $-Y$ et alternativement sur les voies 1 et 2 pour la visualisation ALT. Cette dernière possibilité s'avère particulièrement intéressante lorsque les fréquences des signaux ne sont pas les mêmes. En revanche, ce mode est à exclure pour mesurer des déphasages (il est alors sous-entendu que les deux signaux ont la même fréquence !) car sur une même verticale, les points de la voie 1 et ceux de la voie 2 ne correspondent pas au même instant du fait d'un déclenchement alternativement sur chaque voie.
2. le type de couplage entre la source de déclenchement et le circuit de déclenchement : il s'agit de filtrer le signal de déclenchement. On gardera toutes les composantes du signal dont les fréquences sont comprises entre 0 et 40 MHz pour le couplage DC, uniquement les composantes alternatives de fréquence supérieure à 10 Hz en couplage AC, uniquement les fréquences inférieures à 10 kHz en couplage LF et uniquement les fréquences supérieures à 10 kHz en couplage HF.

e) Réglage du déclenchement

Un certain nombre de paramètres du déclenchement peuvent être choisis :

- le niveau ou « LEVEL » : on choisit la tension seuil du déclenchement. Quand le signal de la source choisie atteint cette valeur, le spot se déclenche. Si l'amplitude du signal de déclenchement est trop petite, aucun déclenchement ne se produira et on n'aura aucune trace sur l'écran. Si on choisit le mode de déclenchement « AUTO », la tension de seuil du déclenchement est limitée aux valeurs prises par le signal : on aura donc toujours une trace sur l'écran.
- la pente : on choisit entre le front montant et le front descendant. Le déclenchement a lieu respectivement lorsque le signal de la source augmente ou diminue.

D'autre part, si l'amplitude du signal est trop faible, le signal peut ne pas être synchronisé. Dans ce cas, il est parfois possible qu'une simple augmentation de la sensibilité verticale suffise à synchroniser.

4. Masse et Terre

La « masse » et la « Terre » sont deux notions différentes que ce paragraphe va préciser.

4.1 Définitions

Par définition, le potentiel électrique est une grandeur définie à une constante près. Il faut donc fixer une origine des potentiels. Cela revient à choisir un point du circuit auquel on affecte *par convention* le potentiel zéro. Ce point sera appelé « masse » du circuit.

La Terre est par définition le conducteur formé de la planète Terre. Lors de l'étude des conducteurs (en deuxième année MP), on verra que le potentiel à la surface externe d'un conducteur est indépendant du point considéré sur la surface. Il est donc possible de choisir la surface de la Terre comme masse, c'est-à-dire comme référence de potentiels. Ce choix n'a cependant rien d'obligatoire, il s'agit d'un choix parmi d'autres.

Les symboles respectifs de la masse et de la Terre sont :



Figure 20.9 Symbole des masse et Terre.

4.2 Importance de la Terre pour la sécurité

Par ailleurs, la Terre joue un rôle électrique primordial : on doit relier les carcasses de tous les appareils électriques à la Terre. Cette connexion évite qu'un utilisateur touchant la carcasse de l'appareil ne serve de conducteur et ne soit traversé par un courant. L'homme est, en effet, conducteur de l'électricité et il existe un danger pour sa santé et sa vie si un courant trop important le traverse. Lorsque la carcasse d'un appareil est reliée à la Terre, un courant qui passerait dans la carcasse de l'appareil rejoindrait directement la Terre plutôt que de traverser un utilisateur touchant cette carcasse avant de rejoindre la Terre. De ce fait, les règles élémentaires de sécurité imposent de relier les carcasses de tous les appareils électriques à la Terre.

4.3 Risque de court-circuit avec la Terre

Lorsque cette règle de sécurité est appliquée, tous les appareils sont reliés à la Terre. Il est souvent commode alors pour le constructeur de relier la Terre au point de référence de potentiel donc à la masse.

Il est alors nécessaire que ce point de référence ne se retrouve pas en différents points du circuit afin d'éviter de court-circuiter un ou plusieurs composants ou appareils.

On reliera dans ce cas toutes les masses entre elles et à la Terre. C'est ce qui explique les confusions entre les deux notions.

4.4 Masse et Terre en pratique

En travaux pratiques, il sera donc nécessaire de gérer concrètement le problème de masse et de Terre.

La première question à résoudre sera de déterminer si la masse d'un appareil est reliée à la Terre. Si ce n'est pas le cas, l'appareil sera dit à *masse flottante*. Une manière de le savoir consiste à mesurer la résistance entre la Terre présente sur les lignes sèches (c'est-à-dire des lignes reliées au secteur sans passer par l'intermédiaire de prises) et la masse de l'appareil. La Terre se trouve à la borne jaune et verte des lignes sèches. En effet, un ohmmètre mesure la valeur de la résistance en envoyant un courant d'intensité connue dans le dipôle, en mesurant la tension aux bornes de celui-ci et en calculant le rapport entre les deux. Si la masse de l'appareil se trouve reliée à la Terre, la différence de potentiel entre la borne Terre des lignes sèches et la borne masse de l'appareil est nulle donc l'ohmmètre donne une valeur quasi-nulle pour la résistance. Dans ce cas, on en déduira que la masse est reliée à la Terre. Sinon on mesure une résistance très grande, proche de la valeur de l'impédance d'entrée de l'appareil de mesure. On conclut dans ce cas que la masse de l'appareil n'est pas reliée à la Terre.

Si la masse de l'appareil n'est pas reliée à la Terre, on peut la choisir en n'importe quel point du circuit sans risque de court-circuit.

Sinon il faudra relier « toutes les Terres » entre elles afin de ne pas avoir de court-circuit.

4.5 Exemple : visualisation de caractéristiques à l'oscilloscope

Pour obtenir une caractéristique, c'est-à-dire la courbe traduisant la relation entre l'intensité du courant traversant le dipôle et la tension à ses bornes, on a besoin de mesurer à la fois ces deux grandeurs. L'intensité sera obtenue expérimentalement en considérant la tension aux bornes d'une résistance placée en série avec le dipôle. On obtient ainsi le montage suivant :

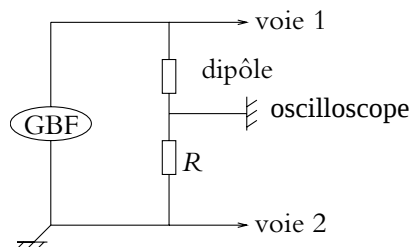


Figure 20.10 Mesure de caractéristiques avec appareil(s) à masse flottante.

Il faut effectivement pouvoir placer la masse de l'oscilloscope entre le dipôle et la résistance. Ceci ne sera possible qu'avec un oscilloscope ou un générateur basse fréquence à masse flottante, c'est-à-dire non reliée à la Terre. Sinon la résistance (dans le cas du branchement de la figure 20.10) ou le dipôle est court-circuité.

Lorsque les masses de l'oscilloscope et du GBF sont reliées à la Terre, on peut s'affranchir de cette difficulté en utilisant le montage suivant :

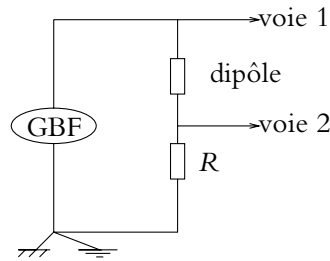


Figure 20.11 Mesure de caractéristiques sans appareils à masse flottante.

Il faudra en outre pouvoir négliger la tension aux bornes de la résistance devant celle aux bornes du dipôle, de manière à pouvoir assimiler la tension visualisée sur la voie 1 à la tension aux bornes du dipôle. Cela nécessite notamment de prendre une faible valeur de résistance, mais ce n'est pas toujours possible (Cf. cas de la diode).

Une autre solution envisageable est l'utilisation d'un transformateur d'isolement :

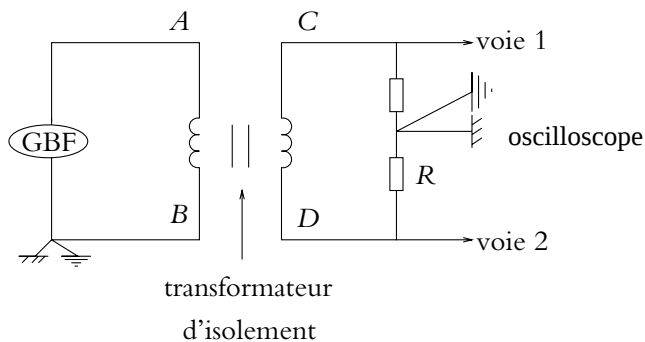


Figure 20.12 Mesure de caractéristiques avec un transformateur d'isolement.

Le transformateur d'isolement permet d'obtenir aux bornes CD du secondaire une tension proportionnelle à celle appliquée au primaire (bornes AB) dans le cas où on travaille avec des signaux sinusoïdaux sans pour autant imposer une masse reliée à la Terre. On peut alors placer celle de l'oscilloscope où on veut, soit ici entre le dipôle et la résistance. Il sera vu dans le cours de deuxième année PSI qu'un transformateur coupe la composante continue. On n'utilisera donc pas un transformateur avec des signaux créneaux dont l'allure en sortie du secondaire sera fortement déformée par le transformateur.

5. Multimètres

5.1 Valeurs moyenne et efficace

On définit la valeur moyenne $\langle x \rangle$ (respectivement la valeur efficace X_{eff}) de $x(t)$ de période T par :

$$\langle x \rangle = \frac{1}{T} \int_{t_0}^{t_0+T} x(t) dt \quad \text{et} \quad X_{\text{eff}} = \sqrt{\frac{1}{T} \int_{t_0}^{t_0+T} x^2(t) dt}$$

On se reportera au chapitre sur la puissance pour une interprétation physique de ces grandeurs.

5.2 Description des multimètres

Un multimètre possède trois bornes d'entrée : une borne commune à tout type de mesure qu'on utilise généralement comme masse, une borne permettant la mesure des intensités (les bornes mA et 10 A) et une borne permettant celle des tensions et des résistances (la borne V Ω).

Suivant son utilisation, un multimètre n'est pas relié au circuit de la même manière :

- en ampèremètre, on l'insère en série dans la branche dont on veut mesurer l'intensité du courant ;
- en voltmètre, on le place en parallèle du dipôle aux bornes duquel on veut mesurer la tension.

Pour ne pas trop perturber les mesures (Cf. exercice A2), il présente donc une résistance de faible valeur quand on l'utilise en ampèremètre et une résistance de grande valeur lorsqu'il est employé comme voltmètre.

En ohmmètre, le multimètre envoie dans le dipôle considéré un courant de faible intensité (dont il connaît la valeur) et mesure la tension qui apparaît de ce fait à ses bornes. Il indique ensuite le rapport de ces deux grandeurs. Ceci explique qu'il est fondamental de sortir un dipôle du circuit dans lequel il se trouve quand on veut mesurer sa résistance. Sinon on mesure la résistance de la totalité du circuit vu entre les deux points considérés.

5.3 Multimètres RMS et non RMS

On distingue deux types de multimètres : ceux qui donnent la valeur efficace vraie (définie plus loin) des signaux et ceux qui ne la donnent pas. Les premiers sont dits RMS.

a) Mesure de la valeur moyenne

Qu'ils soient RMS ou non, tous les multimètres fournissent une mesure de la valeur moyenne de la tension ou de l'intensité lorsqu'on sélectionne la bonne fonction (mode DC).

b) Mesures des valeurs efficaces

Lorsqu'on a effectivement un multimètre RMS, il est possible d'obtenir en sélectionnant la bonne fonction :

- la valeur efficace de la composante alternative,
- la valeur efficace vraie (à savoir du signal complet en tenant compte de la composante continue et de la composante alternative).

Lorsqu'on enclenche le bouton « AC/DC », on garde la totalité du signal (composante continue et composante alternative) alors que non enclenché, on supprime la composante continue : l'entrée du multimètre comporte un condensateur en amont de l'appareil de mesure proprement dit. Il s'agit de la même opération que lorsqu'on passe du mode DC au mode AC d'un oscilloscope.

Si le multimètre n'est pas RMS, il est possible qu'il fournisse comme valeur efficace le produit de la valeur moyenne X_{moy} du signal redressé par 1,1. Dans ce cas, ce sera effectivement la valeur efficace X_{eff} uniquement pour les signaux sinusoïdaux centrés. En effet, le coefficient 1,1 correspond au rapport¹ de la valeur moyenne X_{moy} par la valeur efficace X'_{eff} pour le signal redressé² correspondant. La valeur efficace X'_{eff} du signal correspondant est égale à la valeur efficace X_{eff} du signal initial puisqu'on calcule la moyenne du carré soit pour un signal sinusoïdal : $\frac{X}{\sqrt{2}}$. Le calcul de la valeur moyenne X'_{moy} d'une sinusoïde redressée revient à calculer la valeur moyenne de la valeur absolue de la sinusoïde qui vaut : $\frac{\pi X}{2}$. Le rapport précédent vaut donc $\frac{\pi}{2\sqrt{2}} \simeq 1,1$.

c) Autres mesures possibles

Il est également possible :

- de mesurer des résistances en enclenchant le bouton Ω (fonctionnement en ohmmètre),
- d'effectuer des mesures en décibels en appuyant sur le bouton « dB » si cette fonction existe.

5.4 Impédance d'entrée des multimètres

L'impédance d'entrée d'un multimètre utilisé en voltmètre est une résistance qui va de 1 M Ω à 14 M Ω suivant la qualité du multimètre.

L'influence sur les mesures de l'impédance d'entrée du multimètre étant de même nature que celle de l'impédance d'entrée d'un oscilloscope, on se reportera au paragraphe correspondant sur l'oscilloscope. Dans la plupart des utilisations, cette impédance sera très grande devant celle du circuit étudié et on pourra négliger le courant

¹ Ce rapport est appelé facteur de forme du signal.

² Voir chapitre sur les diodes.

traversant le multimètre. Il ne faudra cependant pas oublier son existence lorsqu'on utilisera des composants dont l'impédance n'est plus négligeable devant l'impédance d'entrée.

L'impédance d'entrée d'un multimètre utilisé en ampèremètre varie suivant le calibre utilisé. On donne en général la chute de tension aux bornes de l'ampèremètre pour un courant donné, par exemple à mi-calibre. Si le constructeur indique $U = 100 \text{ mV}$ sur les calibres de $200 \mu\text{A}$ et 200 mA , à mi-calibre, la résistance interne de l'ampèremètre est donc de 1Ω sur le calibre $200 \mu\text{A}$ et de $1 \text{ k}\Omega$ sur le calibre $200 \mu\text{A}$.

5.5 Bande passante des multimètres

L'un des inconvénients majeurs des multimètres est leur bande passante relativement faible. Ces appareils se comportent comme des filtres passe-bas qui ne laissent donc passer que les basses fréquences. Ils ne donnent des mesures correctes que dans cette bande de fréquences ; au-delà, les mesures effectuées ne seront plus valables. La fréquence de coupure est souvent de l'ordre de quelques dizaines de kHz.

On notera que la bande passante des oscilloscopes est beaucoup plus grande, de l'ordre de quelques MHz. On aura donc recours si nécessaire plutôt à l'oscilloscope pour tracer expérimentalement un diagramme de Bode en gain.

6. Modèles réels des dipôles linéaires passifs

Les modèles adoptés au cours des chapitres antérieurs pour les résistors, les bobines et les condensateurs sont des modèles parfaits au sens où ils ne sont pas toujours valables en pratique. Par exemple, quand la fréquence devient très importante, les comportements observés de ces composants ne correspondent plus à ce que leur modèle parfait décrit.

L'objet de ce paragraphe est donc de décrire les modèles des composants réels. On utilisera pour cela les trois éléments parfaits dont le comportement est décrit par une relation entre l'intensité i du courant qui les traverse et la tension à leurs bornes :

- résistor de résistance $u = Ri$,
- bobine d'inductance $u = L \frac{di}{dt}$,
- condensateur de capacité $i = C \frac{du}{dt}$.

6.1 Modèles du condensateur

La caractéristique principale d'un condensateur est sa capacité électrostatique C_0 . Cependant on observe les différents points suivants.

- En régime continu, un condensateur parfait chargé conserve la même charge au cours du temps. Expérimentalement, un condensateur réel chargé se décharge progressivement au cours du temps. Il s'agit d'un phénomène de fuites liées à une conductivité extrêmement faible mais non nulle de l'isolant (qui n'est pas parfait) compris entre les deux armatures.
- Il y a dissipation d'énergie électrique au niveau du diélectrique et des connexions lorsque la tension est alternative.
- Le comportement évolue en fonction de la fréquence :
 - à faible fréquence, on observe des effets capacitifs,
 - à haute fréquence, ce seront des effets inductifs,
 - au voisinage d'une fréquence ω_0 , les effets observés sont résistifs.

Pour rendre compte de ces observations, il est possible d'adopter les modèles suivants :

- **Modèle très basses fréquences :**

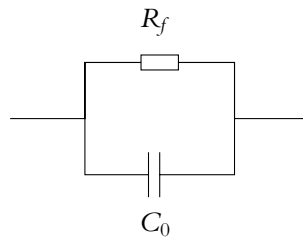


Figure 20.13 Modèle basses fréquences d'un condensateur.

On se trouve en régime quasi-continu : on ne tient compte que des phénomènes de fuite en plaçant une résistance R_f , dite résistance de fuite, en parallèle avec la capacité électrostatique C_0 . L'ordre de grandeur de R_f varie du $M\Omega$ à quelques centaines de $M\Omega$.

- **Modèle large bande :**

L'objectif est de rendre compte du comportement du condensateur sur une large plage de fréquences, ce qui donne le nom de ce modèle. On utilise pour ce faire un circuit R, L, C série :

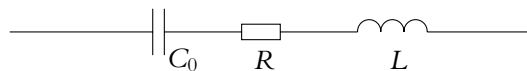


Figure 20.14 Modèle large bande d'un condensateur.

Il s'agit de rendre compte à haute fréquence des observations, à savoir que des effets inductifs et résistifs se manifestent et qu'il existe des phénomènes de pertes.

Exemple

Pour un condensateur en céramique de capacité électrostatique $C_0 = 10 \text{ nF}$, les valeurs des composants du modèle sont :

$$\begin{cases} R = 20 \, \Omega \\ L = 1,5 \text{ nH} \\ C_0 = 10 \text{ nF} \end{cases}$$

Pour ce type de composants, la fréquence de coupure est de 40 MHz, ce qui signifie que pour une fréquence $f \ll 40 \text{ MHz}$ ce sont les effets capacitifs qui dominent et que pour une fréquence $f \gg 40 \text{ MHz}$ ce sont les effets inductifs.

- **Modèle bande étroite :**

Pour certaines applications, les signaux ont une fréquence fixe ou située dans une bande de fréquences très faible (par exemple en sortie d'un filtre sélectif). Dans ce cas, il est souvent suffisant de connaître le comportement du condensateur dans la bande de fréquences concernée. On adopte alors l'un des modèles suivants :

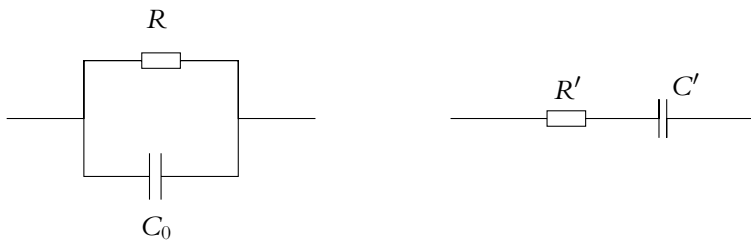


Figure 20.15 Modèles bande étroite d'un condensateur.

On peut établir qu'au voisinage d'une fréquence donnée, il est possible de choisir les valeurs de R' et C' en fonction de celles de R et C_0 (et réciproquement) pour que ces deux modèles soient équivalents.

En général dans les applications envisagées, les effets capacitifs seront prépondérants.

6.2 Modèles d'une bobine réelle

Le comportement d'une bobine dépend beaucoup des conditions dans lesquelles on l'utilise, par exemple de la présence ou non d'un noyau en fer ou en ferrite. Pour les signaux de forte amplitude, il est préférable de considérer directement les caractéristiques magnétiques³. En revanche pour les signaux de faible amplitude, un modèle souvent applicable compte tenu des observations expérimentales est caractérisé essentiellement par une inductance L_0 . Cependant, comme pour le condensateur, cela suppose de négliger certaines observations.

³ Certains de ces aspects seront développés dans le cours de deuxième année de la filière PSI.

- Une bobine présente en régime continu des phénomènes de dissipation correspondant à la résistance du bobinage, à savoir des fils conducteurs qui le constituent.
- En régime alternatif, on observe des phénomènes dissipatifs par effet Joule dans le bobinage ainsi que des phénomènes dissipatifs liés aux pertes d'origine magnétique (que ce soit par hystérésis ou par courants de Foucault⁴).
- En haute fréquence, on constate l'apparition d'effets capacitifs qu'on peut expliquer par l'existence de « mini-condensateurs » formés par les spires entourées d'isolant : deux spires consécutives constituent les armatures et sont séparées par le diélectrique qu'est l'isolant.

On pourra par exemple adopter les modèles suivants :

- **Modèle très basses fréquences :**

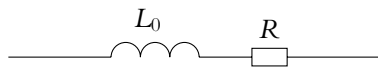


Figure 20.16 Modèle basse fréquence d'une bobine.

On ne tient compte que des phénomènes dissipatifs qui apparaissent à toute fréquence y compris nulle (ce qui correspond au régime continu) en plaçant une résistance R , de l'ordre d'une dizaine d'ohms, en série avec l'inductance L_0 .

- **Modèle large bande :**

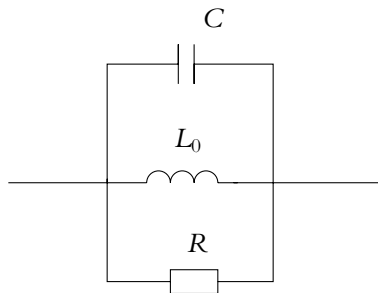


Figure 20.17 Modèle large bande d'une bobine.

Il s'agit de rendre compte de l'évolution du comportement de la bobine dans une large plage de fréquences et notamment, à haute fréquence, des effets capacitifs et des phénomènes dissipatifs modélisés par une résistance.

- **Modèle bande étroite :**

Comme pour le condensateur, on peut être amené à utiliser la bobine avec des signaux à fréquence fixe ou dont la fréquence est comprise dans une bande de

⁴Ces aspects seront étudiés en deuxième année de la filière PSI.

fréquences réduite. Dans ce cas, il est souvent suffisant de connaître le comportement de la bobine dans la bande de fréquences concernée. On adopte alors l'un des modèles suivants :

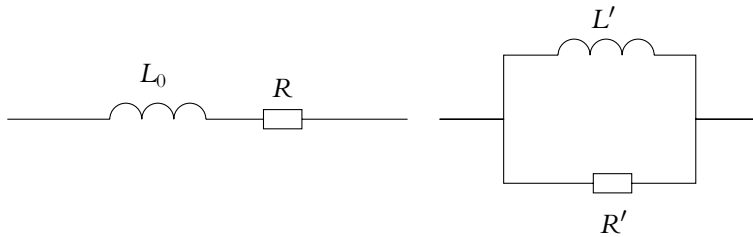


Figure 20.18 Modèles bande étroite d'une bobine.

On peut établir (Cf. exercice B.1. du chapitre sur le régime sinusoïdal) qu'au voisinage d'une fréquence donnée, il est possible de choisir les valeurs de R' et L' en fonction de celles de R et L_0 (et réciproquement) pour que ces deux modèles soient équivalents.

6.3 Modèles d'un résistor

Même les résistors ne sont pas parfaits. À basse fréquence, on observe le comportement décrit par le modèle parfait : une résistance pure.

Cependant, quand la fréquence augmente, deux types de comportements peuvent apparaître essentiellement suivant la technologie du composant :

- des effets capacitifs comme pour les résistances à couches de carbone,
- des effets inductifs comme pour les résistances en spirales.

Les modélisations possibles sont donc :

- **modèle basse fréquence** : résistance pure,
- **modèle large bande** : ce modèle devant être valable dans un vaste domaine de fréquences, on adopte le modèle suivant :

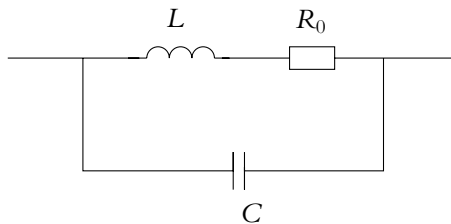


Figure 20.19 Modèle large bande d'une résistance.

- **modèle bande étroite** : on utilisera soit une résistance en série avec une inductance pour des effets inductifs soit une résistance en parallèle avec une capacité pour des effets capacitifs.

6.4 En pratique

Malgré ce qui vient d'être exposé, on aura l'habitude de choisir les modélisations suivantes :

- une capacité pure pour un condensateur,
- une inductance en série avec une résistance pour une bobine,
- une résistance pure pour une résistance.

Si ce modèle n'est pas valable, c'est-à-dire si on constate des différences de comportement notamment en hautes fréquences, il pourra être envisagé d'adopter l'un des modèles proposés plus haut.

On notera cependant que l'adoption d'un modèle bande étroite impose aux valeurs des composants de dépendre de la bande de fréquences considérée.

7. Mesure de déphasages

On rappelle que le déphasage n'est mesurable que pour des signaux de même fréquence. Il n'y a aucun problème de déclenchement à l'oscilloscope : la source peut être indifféremment l'une ou l'autre des voies. On n'utilisera surtout pas le mode de déclenchement ALT : l'origine des temps pour l'une ou l'autre des voies n'étant pas concomitante, on ne peut effectuer de mesures de déphasage.

Il existe plusieurs méthodes pour effectuer ces mesures.

7.1 Méthode de Lissajous en XY

Deux sinusoïdes de même fréquence

$$u_1(t) = U \cos(2\pi ft) \quad u_2(t) = V \cos(2\pi ft + \varphi)$$

visualisées en mode XY dessinent sur l'écran de l'oscilloscope une ellipse dont les caractéristiques dépendent du déphasage φ .

Les points A et B s'obtiennent pour $u_1(t) = 0$ soit $2\pi ft = \frac{\pi}{2} [\pi]$ ou $t = \frac{1}{4f} \left[\frac{1}{2f} \right]$. Les valeurs de $u_2(t)$ correspondantes sont $\pm V \sin \varphi$. La distance AB correspond donc à $2V \sin \varphi$. La distance MN vaut $2V$ puisque M et N correspondent aux valeurs

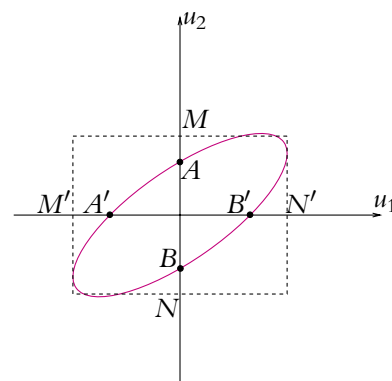


Figure 20.20 Figure de Lissajous pour la mesure d'un déphasage.

extrêmes de $u_2(t)$. On en déduit

$$|\sin \varphi| = \frac{AB}{MN}$$

Par un raisonnement analogue, on peut établir que

$$|\cos \varphi| = \frac{A'B'}{M'N'}$$

En pratique, on commencera par bien centrer l'ellipse : lorsqu'on place l'une des voies sur la position « ground » ou zéro, on obtient un segment parallèle à l'un des axes et on le place au milieu de l'écran. On répète l'opération pour la deuxième voie. On mesure alors AB (ou $A'B'$) puis MN (ou $M'N'$). Cette dernière mesure est facilitée en annulant l'une des voies (c'est-à-dire en la positionnant sur « ground »).

Le signe du déphasage s'obtient en regardant lequel des signaux est en avance sur l'autre.

Cette méthode est très performante pour des déphasages de 0 ou π : l'ellipse est alors une droite. On aura le dédoublement dès que le déphasage ne sera plus rigoureusement égal à 0 ou π . On l'utilisera, par exemple, pour mesurer la pulsation de résonance en intensité d'un circuit R, L, C série pour laquelle on a un déphasage nul.

En revanche, elle est peu précise dans les autres cas. Notamment lorsque le déphasage est proche de $\frac{\pi}{2}$, le sinus varie peu autour de 1 et on aura des difficultés à détecter ces écarts. On évitera donc d'utiliser la méthode de Lissajous sauf pour des déphasages de 0 ou π ou proches de ces valeurs.

7.2 Mesures directes sur l'oscilloscope

On s'intéresse ici aux variations du signal autour de sa valeur moyenne, il faut donc se placer en mode AC.

On règle préalablement soigneusement le zéro des deux signaux et on les fait coïncider.

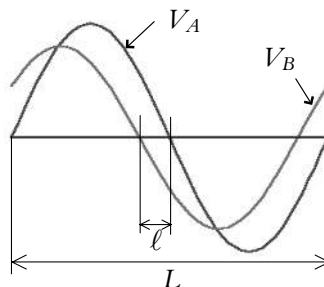


Figure 20.21

On a :

$$\varphi = 360 \times \frac{\ell}{L} \text{ en degrés.}$$

► **Remarque :**

1. Si on décalibre la base de temps pour inscrire une période sur 9 carreaux , $\varphi = 40\ell$ (toujours en degrés),
2. Si on met une demi-période sur 9 carreaux, $\varphi = 20\ell$ (toujours en degrés),
3. Si on veut régler un déphasage de $\pi/4$ radians (pour déterminer une fréquence de coupure par exemple), il peut être intéressant d'adapter cette méthode : en mettant une période sur 8 carreaux, un déphasage de $\pi/4$ est représenté par 1 carreau.

7.3 Utilisation d'un phasemètre numérique

Dans ce cas, il suffit de lire la valeur donnée par l'appareil après avoir effectué les bons réglages. Cette méthode suppose de faire confiance à l'appareil et au manipulateur... Avoir une idée du résultat évite des erreurs grossières !

Il est également possible de faire la mesure directement avec un oscilloscope numérique : dans le menu des mesures, on trouvera la fonction correspondante.

8. Mesures d'impédances par diviseur de tension

8.1 Principe de la mesure

Lors de l'étude des ponts diviseurs de tension, on a vu que la tension aux bornes d'une partie des résistances était égale à une fraction de la tension imposée à la totalité. Ce résultat peut être exploité pour mesurer des impédances.

Le schéma du montage est donc le suivant :

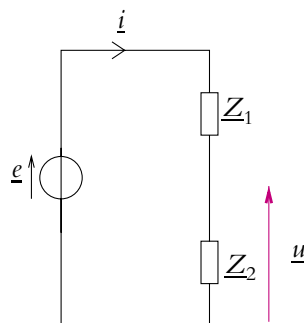


Figure 20.22 Mesure d'impédances par diviseur de tension.

On suppose connue l'impédance $\underline{Z}_1$ et on cherche à déterminer $\underline{Z}_2$. La relation du pont diviseur de tension donne :

$$\underline{u} = \frac{\underline{Z}_2}{\underline{Z}_1 + \underline{Z}_2} \underline{e}$$

On mesure les caractéristiques de la tension $\underline{e}$ en prenant $\underline{Z}_1 = 0$. En mesurant celles de $\underline{u}$ pour une valeur connue de $\underline{Z}_1$, on en déduit la détermination de $\underline{Z}_2$.

8.2 Mesure d'une résistance

C'est le cas le plus simple où les impédances $\underline{Z}_1$ et $\underline{Z}_2$ sont des résistances notées respectivement R_1 et R_2 . Il n'y a alors pas de déphasage et la seule mesure de l'amplitude permet d'obtenir la valeur de R_2 à partir de la relation : $U = \frac{R_2}{R_1 + R_2} E$ soit :

$$R_2 = \frac{R_1 (E - U)}{U}$$

En général, on utilise la méthode de la demi-tension. Il s'agit de prendre une résistance R_1 variable et de modifier sa valeur jusqu'à obtenir une amplitude U égale à la moitié de celle de E . La relation précédente devient alors :

$$R_2 = R_1$$

et il suffit de « lire » la valeur de R_1 pour connaître celle de R_2 cherchée.

On utilise cette méthode pour mesurer la résistance d'un résistor bien sûr mais également les résistances d'entrée ou de sortie d'un montage comme par exemple la résistance d'entrée d'un oscilloscope ou la résistance de sortie d'un GBF.

8.3 Mesure d'une inductance ou d'une capacité

Dans ces deux cas, les impédances sont effectivement complexes et non réelles : il est donc nécessaire de mesurer les amplitudes E et U de $\underline{e}$ et $\underline{u}$ ainsi que le déphasage φ entre les deux tensions $\underline{e}$ et $\underline{u}$ à une fréquence donnée.

On obtient alors la relation en notation complexe :

$$\underline{Z}_2 = \frac{\underline{Z}_1 (\underline{e} - \underline{u})}{\underline{u}}$$

On en déduit alors l'amplitude de $\underline{Z}_2$:

$$|\underline{Z}_2| = \frac{|\underline{Z}_1| |E - U \exp(j\varphi)|}{U}$$

et son argument par :

$$\text{Arg} \underline{Z}_2 = \text{Arg} \underline{Z}_1 + \text{Arg} (E - U \exp(j\varphi)) - \varphi$$

C'est la méthode utilisée au paragraphe a) pour déterminer l'impédance d'entrée de l'oscilloscope. L'impédance connue $\underline{Z}_1$ est une résistance réglable R_1 . L'impédance inconnue $\underline{Z}_2$ est l'impédance d'entrée de l'oscilloscope, équivalente à une résistance R_o en parallèle avec un condensateur de capacité C_o . Elle est égale à $\underline{Z}_o = \frac{R_o}{1 + jR_o C_o \omega}$. Dans un premier temps, on travaille à très basse fréquence (100 Hz environ) de telle sorte que $R_o \ll \frac{1}{C_o \omega}$ pour déterminer R_o . Puis, à fréquence plus élevée (10 kHz par exemple), on mesure

$$\left| \frac{u}{\underline{e}} \right| = \left| \frac{1}{1 + \frac{\underline{Z}_2}{\underline{Z}_1}} \right| = \frac{1}{\sqrt{\left(1 + \frac{R_1}{R_o}\right)^2 + (C_o R_1 \omega)^2}}$$

ce qui permet de déterminer la valeur de C_o .

A. Applications directes du cours

1. Résistance interne d'un ampèremètre ou d'un voltmètre

1. Pour réaliser de bonnes mesures, un ampèremètre doit-il avoir une résistance interne forte ou faible ? Par rapport à quelle grandeur définit-on cette « force » ?
2. Mêmes questions pour un voltmètre.

2. Comment utiliser un voltmètre et un ampèremètre pour mesurer une résistance ?

On étudie l'influence de la valeur de la résistance des appareils de mesure lors de la mesure d'une résistance avec un ampèremètre et un voltmètre. On effectue cette détermination avec l'un des trois montages suivants en appliquant la loi d'Ohm aux valeurs lues sur le voltmètre et l'ampèremètre.

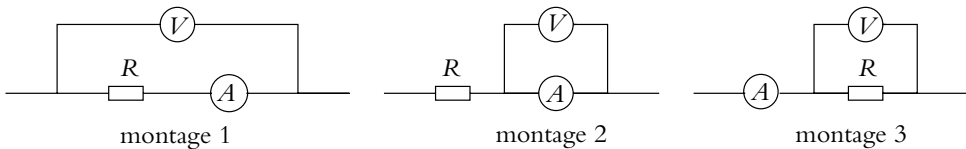


Figure 20.23

1. Quel(s) montage(s) sont-ils effectivement utilisables pour la mesure de la résistance ?
2. Pour chacun des montages utilisables, exprimer R_{mes} , valeur estimée de la résistance, en fonction de la valeur vraie R et des résistances internes R_A et R_V respectivement de l'ampèremètre et du voltmètre.
3. En fonction de la valeur de R à mesurer, comment choisir le montage le plus adapté ?

B. Exercices et problèmes

1. Résistance de fuite d'un condensateur

Soit un condensateur de capacité C et présentant une résistance de fuite R_f . On peut modéliser le condensateur par l'association en parallèle de R_f et C . On mesure la tension aux bornes du condensateur à l'aide d'un voltmètre électronique parfait (de résistance interne infinie). Ce condensateur ayant été chargé sous une tension E à l'aide d'une source idéale de tension, on ouvre le circuit. Au bout d'un temps T , on constate que la tension indiquée par le voltmètre n'est plus que de $E' < E$.

1. Comment peut-on expliquer ces observations ?
2. Donner l'expression de R_f en fonction de C , E , E' et T .

2. Étude expérimentale d'un circuit RC, d'après ENSTIM

On considère le circuit suivant comprenant une résistance de valeur R , un condensateur de capacité C et une alimentation stabilisée de tension à vide E :

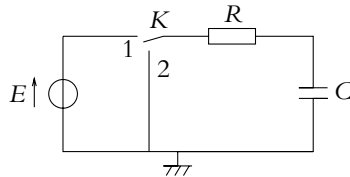


Figure 20.24

1. À l'instant $t' = 0$, on place l'interrupteur K en position 1, le condensateur est déchargé. Décrire ce qui se passe.
2. On suppose que le régime permanent a été atteint. On place l'interrupteur K en position 2 à $t = 0$. Quelle est la nature de la courbe représentant $\ln\left(\frac{E}{u_C}\right)$ en fonction du temps ?

On réalise l'expérience avec les valeurs $R = 10 \text{ M}\Omega$, $C = 10 \text{ }\mu\text{F}$ et $E = 10 \text{ V}$. On branche un voltmètre numérique (calibre 20 V) aux bornes du condensateur et on étudie la décharge du condensateur à partir de l'instant de date $t = 0$ où on place l'interrupteur en position 2. On relève les valeurs suivantes :

$t(\text{s})$	5	10	15	20	30	45	60	90	120	150
$u_C (\text{V})$	9,05	8,19	7,41	6,70	5,49	4,07	3,01	1,65	0,91	0,50

3. L'expérience est-elle en accord avec la théorie ?
4. Si oui, justifier. Si non, pour quel composant a-t-on utilisé un modèle inadapté aux caractéristiques du circuit ? On indiquera alors quel serait le modèle adapté et les valeurs des caractéristiques de ce dernier compte tenu des données expérimentales.
5. Quelle aurait été la condition à vérifier pour avoir accord parfait entre théorie et expérience ?

On souhaite visualiser sur un oscilloscope en mode DC les courbes théoriques de charge et de décharge du condensateur. On choisit comme valeurs $R = 10 \text{ k}\Omega$ et $C = 10 \text{ nF}$. L'alimentation stabilisée est remplacée par un générateur basse fréquence (GBF) et on place l'oscilloscope aux bornes du condensateur.

6. Rappeler le modèle électrique du GBF et donner les valeurs des éléments du modèle.
7. Même question pour l'oscilloscope.
8. Montrer qu'on peut se ramener au cas théorique tout en tenant compte des paramètres du GBF et de l'oscilloscope. On précisera les nouvelles valeurs de la tension, de la résistance et du condensateur.
9. Déterminer les conditions qui permettent de n'avoir pas à tenir compte des paramètres du GBF et de l'oscilloscope. En déduire la justification du choix des valeurs employées pour R et C .

10. Quel signal du GBF doit-on choisir à la sortie du GBF pour observer sur l'oscilloscope la charge et la décharge du condensateur ?
11. Comment doit-on choisir la fréquence de ce signal pour observer pratiquement toute la charge puis de la décharge du condensateur ? Dans ce cas, représenter l'allure du signal obtenu.

3. Étude expérimentale d'un circuit RC , d'après ENSAIT 1999 concours A

Soit un circuit comprenant en série un condensateur de capacité C variable, une bobine d'inductance L et de résistance R , un générateur idéal de tension basse fréquence d'amplitude E et de fréquence f fixe et un ampèremètre de résistance négligeable (sauf indication contraire).

1. Représenter le schéma du circuit.
2. Déterminer l'expression de l'intensité parcourant le circuit.
3. Établir qu'en faisant varier la valeur de la capacité, l'intensité passe par un maximum I_0 . On notera C_0 la valeur de la capacité correspondante.
4. On détermine les valeurs C_1 et C_2 de la capacité correspondant à une intensité $\frac{I_0}{\sqrt{2}}$. Expliquer les expressions de C_1 et C_2 en fonction de L , R et ω la pulsation.
5. On fait l'hypothèse que C_1 et C_2 sont peu différentes de C_0 . Expliciter l'approximation que cela permet de faire.
6. Montrer que sous cette hypothèse les valeurs C_1 et C_2 sont équidistantes de C_0 .
7. Que peut-on dire du déphasage pour $C = C_0$?
8. Même question pour C_1 et C_2 . Comment peut-on qualifier ces valeurs ?
9. Quelle est la meilleure façon de déterminer expérimentalement C_0 ?
10. Pourquoi est-il plus avantageux ici de mesurer C_1 et C_2 que C_0 ?
11. Définir le facteur de surtension Q aux bornes de la capacité et montrer qu'il correspond également au facteur de surtension aux bornes de l'inductance.
12. Exprimer Q en fonction de C_1 et C_2 .
13. Dédurre de ce qui précède les expressions de L et R en fonction de C_1 , C_2 et de la fréquence f .
14. Que se passe-t-il si on ne néglige plus la résistance R_A de l'ampèremètre ? Préciser.
15. Application numérique : On donne $C_1 = 120 \text{ nF}$, $C_2 = 130 \text{ nF}$ et $f = 100 \text{ kHz}$. Déterminer le facteur de surtension et les valeurs de R et L en négligeant R_A puis lorsque $R_A = 0,1 \Omega$.
16. Les approximations faites sont-elles justifiées ?

21

T.P. Cours : Amplificateur opérationnel

On s'intéresse dans ce chapitre à un composant très fréquemment utilisé : l'amplificateur opérationnel. On décrit le modèle idéal de ce composant à partir des observations expérimentales et de quelques données numériques. On précise ensuite quelques écarts par rapport à ce modèle observables expérimentalement et leur influence sur les résultats expérimentaux. On étudie ensuite quelques montages utilisant ce composant aussi bien en régime linéaire que saturé.

1. Modèle de l'amplificateur idéal

1.1 Description

L'amplificateur opérationnel est un composant intégré qui doit être alimenté par une tension continue $-15\text{ V} / +15\text{ V}$.

Il possède deux bornes d'entrée notées E^+ et E^- appelées respectivement *entrée non inverseuse* et *entrée inverseuse*. On verra ultérieurement les raisons de cette dénomination.

Il présente enfin une borne de sortie souvent notée S .

Habituellement on emploie le schéma des figures 21.1 et 21.2.

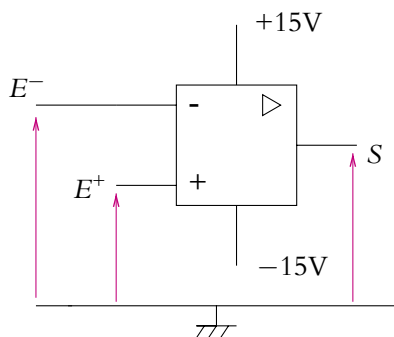


Figure 21.1 Représentation de l'amplificateur opérationnel.

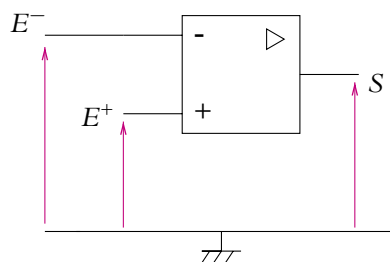


Figure 21.2 Représentation simplifiée de l'amplificateur opérationnel.

On notera que sur la figure 21.1, on a fait apparaître les fils de connexion de l'alimentation. Pour ne pas surcharger les schémas, on adoptera celui de la figure 21.2 où l'alimentation n'apparaît pas. Il ne faudra cependant jamais oublier que, sans elle, le composant n'a pas les propriétés qui vont être décrites par la suite.

1.2 Régimes de l'amplificateur opérationnel

a) Caractéristique de transfert statique (PCSI, PTSI)

On note ϵ la différence de potentiel imposée entre les deux bornes d'entrée :

$$\epsilon = v_+ - v_-$$

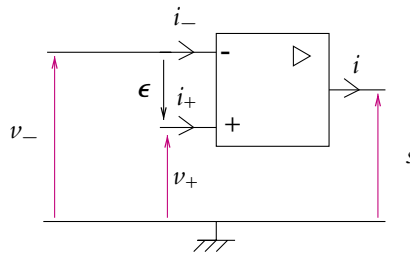


Figure 21.3 Notations associées à l'amplificateur opérationnel.

Pour étudier la réponse de la tension de sortie s en fonction de la tension différentielle d'entrée ϵ , on peut, par exemple, utiliser le montage de la figure 21.4 en fixant la tension de l'entrée inverseuse à 0 et en connectant un générateur sinusoïdal à l'entrée non inverseuse. Dans ce cas, $\epsilon = v_+ - v_- = v_+$. En visualisant en mode XY les tensions ϵ et s , on obtient les résultats suivants :

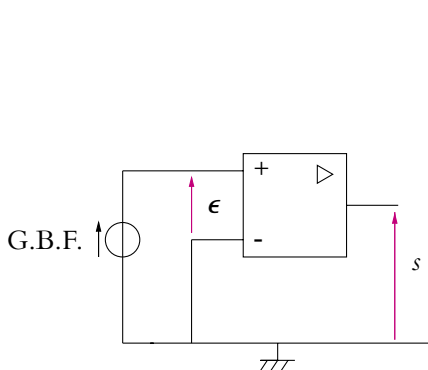


Figure 21.4 Montage pour tracer la caractéristique de transfert.

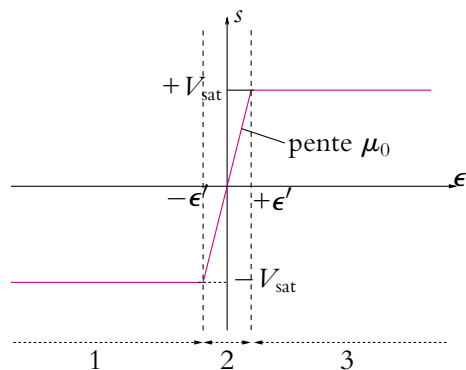


Figure 21.5 Caractéristique de transfert statique de l'amplificateur opérationnel.

Cette caractéristique de transfert fait apparaître trois zones de fonctionnement numérotées de 1 à 3 et correspondant à deux types de régime.

b) Régime saturé

Le *régime saturé* correspond aux zones 1 et 3. Dans ce cas, la tension de sortie de l'amplificateur opérationnel est constante et indépendante de la valeur de la tension différentielle d'entrée ϵ . On dit que l'amplificateur opérationnel sature.

La tension de saturation est notée V_{sat} et sa valeur est habituellement comprise entre 12,5 et 14,5 V pour une alimentation de l'amplificateur opérationnel de $-15 \text{ V} / +15 \text{ V}$. En pratique, la tension de saturation négative n'est pas forcément égale en valeur absolue à la tension de saturation positive. La valeur de cette tension dépend de l'alimentation du composant. On verra ultérieurement qu'il s'agit d'un fonctionnement en régime non linéaire.

c) Régime linéaire

Le *régime linéaire* correspond à la zone 2 et est caractérisé par la relation linéaire :

$$s = A_d \epsilon = A_d (v_+ - v_-)$$

en notant A_d le coefficient d'amplification différentielle en régime continu. Sa valeur typique est de l'ordre 10^5 .

On en déduit que la zone de passage de $-V_{\text{sat}}$ à $+V_{\text{sat}}$ est très étroite. En effet, l'écart sur les valeurs de ϵ correspondant au régime linéaire est :

$$\Delta \epsilon_{\text{max}} \simeq 2 \times \frac{13}{10^5} = 26 \cdot 10^{-5} \text{ V}$$

Comme $26 \cdot 10^{-5} \ll 1$, il sera donc en pratique délicat de se placer dans cet intervalle de valeurs pour la tension ϵ : le fonctionnement en régime linéaire ne doit pas être considéré comme le seul fonctionnement de l'amplificateur opérationnel. Au contraire, il s'agit *a priori* d'une situation beaucoup plus exceptionnelle que le fonctionnement en saturation. Cela n'empêchera pas dans nombre de situations expérimentales, d'exercices et de problèmes de se placer dans l'hypothèse d'un fonctionnement en régime linéaire. Il s'agit en effet d'un fonctionnement malgré tout fréquent : l'amplificateur opérationnel est rarement utilisé tout seul. En particulier, une liaison entre l'entrée inverseuse et la sortie peut permettre un fonctionnement en régime linéaire comme on le verra ultérieurement.

1.3 Modèle de l'amplificateur opérationnel idéal

Très souvent, on utilise une idéalisation de ce composant du fait des valeurs de certaines grandeurs caractéristiques.

Le modèle idéal de l'amplificateur opérationnel¹ correspond à faire les hypothèses suivantes :

- absence de courants d'entrée :

$$i_+ = i_- = 0$$

parfaitement justifiée par leur ordre de grandeur qui est de l'ordre de 30 pA ($30 \cdot 10^{-12}$ A) pour la génération des TL081 et de l'ordre de 80 nA pour celle des 741,

- gain différentiel A_d infini :

$$A_d \rightarrow +\infty$$

justifié aussi par sa valeur typique de $2 \cdot 10^5$ aussi bien pour les TL081 que les 741,

- source idéale de tension en sortie, c'est-à-dire absence de résistance interne en sortie.

La caractéristique statique de l'amplificateur idéal est donc :

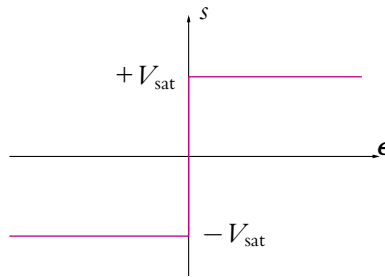


Figure 21.6 Caractéristique de transfert statique d'un amplificateur opérationnel idéal.

Sur le symbole, le triangle sera suivi du signe « ∞ » pour un amplificateur opérationnel idéal.

1.4 Conséquence du caractère idéal sur le régime linéaire

Le modèle précédent ne présage en rien du type de régime dans lequel se trouve l'amplificateur opérationnel. Il peut être soit en régime linéaire soit en régime saturé. S'il est en régime linéaire, l'hypothèse selon laquelle $A_d \rightarrow +\infty$ implique que la zone de fonctionnement linéaire de la caractéristique se réduit à une partie verticale et impose donc que :

$$\epsilon = 0$$

¹ Les amplificateurs opérationnels qui sont généralement utilisés en travaux pratiques sont les TL081 et les 741 pour lesquels on donnera par la suite les valeurs des grandeurs caractéristiques.

Il faut bien noter que cette condition résulte de deux hypothèses totalement indépendantes :

- amplificateur opérationnel idéal,
- fonctionnement en régime linéaire.

Ce qui vient d'être dit peut se résumer par le tableau suivant :

	Amplificateur opérationnel idéal	Amplificateur opérationnel non idéal
régime linéaire	$i_+ = i_- = 0$ $A_d \rightarrow +\infty$ $\epsilon = 0$	$s = A_d \epsilon$
régime saturé	$i_+ = i_- = 0$ $A_d \rightarrow +\infty$ $V_S = \pm V_{\text{sat}}$	$V_S = \pm V_{\text{sat}}$

2. Deux montages simples

Dans ce paragraphe, on présente deux montages à amplificateur opérationnel, le *sui-veur* et l'*amplificateur non-inverseur*, qui seront étudiés en détail plus loin dans le chapitre. Ils pourront servir de montage permettant de mettre en évidence les phénomènes présentés dans le paragraphe 3.

2.1 Montage suiveur

a) Montage - Étude expérimentale

Il s'agit du montage donné par la figure 21.7.

Le générateur est un générateur réel qu'on représente par son modèle de Thévenin à savoir une source idéale de tension de f.e.m. $\underline{E}_T$ en série avec une impédance $\underline{Z}_T$.

Si on choisit une tension d'alimentation sinusoïdale, d'amplitude de l'ordre de quelques volts et de fréquence de l'ordre de quelques kilohertz, on observe une tension de sortie identique à la tension d'entrée.

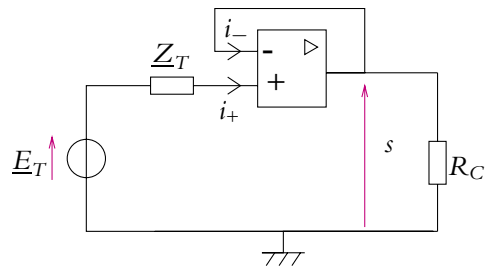


Figure 21.7 Montage suiveur.

b) Étude dans l'hypothèse d'un amplificateur opérationnel idéal

On suppose l'amplificateur opérationnel idéal donc $i_+ = i_- = 0$ A. On constate expérimentalement qu'il fonctionne en régime linéaire donc $v_+ = v_-$. On peut écrire :

$$\epsilon = v_+ - v_- = 0 \quad \text{soit} \quad \underline{s} = \underline{E}_T$$

On peut alors justifier le nom donné à ce montage : la tension de sortie est égale à la tension à vide du générateur placé à l'entrée du montage, la sortie « suit » l'entrée.

2.2 Montage amplificateur non inverseur

a) Montage - Étude expérimentale

Il s'agit du montage de la figure 21.8.

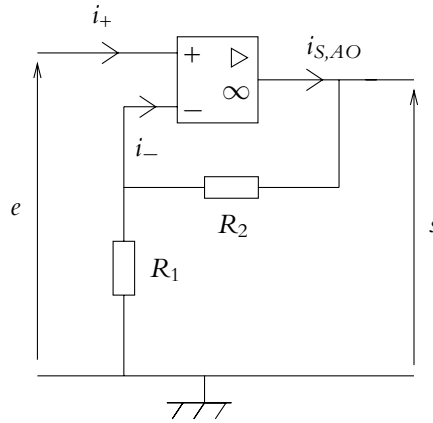


Figure 21.8 Montage non inverseur.

Si on choisit une tension d'alimentation sinusoïdale, d'amplitude de l'ordre de quelques dixièmes de volts et de fréquence de l'ordre de quelques kilohertz, on observe que la tension de sortie est en phase avec la tension d'entrée, le gain en tension étant égal à $1 + \frac{R_2}{R_1}$.

b) Étude dans l'hypothèse d'un amplificateur opérationnel idéal

On constate expérimentalement un fonctionnement en régime linéaire. Si on suppose l'amplificateur opérationnel idéal, on a donc :

$$\epsilon = v_+ - v_- = 0 \quad (21.1)$$

Or on a $v_+ = e$ et, par un pont diviseur de tension puisque $i_- = 0$, $v_- = \frac{R_1}{R_1 + R_2}s = \alpha s$ en notant $\alpha = \frac{R_1}{R_1 + R_2}$.

En reportant dans (21.1), on en déduit :

$$\alpha s = e$$

soit

$$\frac{s}{e} = \frac{1}{\alpha} = 1 + \frac{R_2}{R_1}$$

La tension de sortie est égale à la tension d'entrée amplifiée d'un facteur $1 + \frac{R_2}{R_1}$ d'où le nom de montage amplificateur.

De plus, le coefficient d'amplification est positif : on « n'inverse pas » en sortie la tension d'entrée. On qualifie donc ce montage de montage amplificateur non inverseur. On remarquera que la tension d'entrée est reliée à la borne non inverseuse de l'amplificateur opérationnel.

3. Quelques écarts au modèle idéal de l'amplificateur opérationnel

Le modèle de l'amplificateur opérationnel idéal n'est qu'un modèle qui est satisfaisant dans une large mesure. Cependant, dans un certain nombre de cas, on observe des écarts entre les comportements réellement obtenus et ceux qui sont prédits dans le cadre du modèle idéal. Il devient alors nécessaire de modifier le modèle pour tenir compte de ces différences. L'objet de cette partie est de décrire un certain nombre de « défauts » de l'amplificateur opérationnel réel par rapport au modèle idéal et de fournir le modèle « correct » correspondant.

3.1 Saturation en tension

La valeur absolue de la tension de sortie de l'amplificateur opérationnel ne peut dépasser une valeur appelée *tension de saturation* et notée V_{sat} . Lorsque le composant est en régime linéaire, il faudra toujours vérifier que la valeur absolue tension de sortie s ne dépasse pas V_{sat} . Si ce n'est pas le cas, les expressions qui auront été obtenues pour s en fonction des tensions appliquées aux entrées inverseuse et non inverseuse ne seront plus valables.

En fait, la saturation basse et la saturation haute ne sont pas symétriques. Cette dissymétrie peut être due d'une part à une éventuelle dissymétrie de l'alimentation continue qui n'est pas exactement de $+15\text{ V}$ et -15 V et d'autre part à une dissymétrie intrinsèque de l'amplificateur opérationnel. On rappelle que les tensions de saturation sont habituellement comprises entre $12,5$ et $14,5\text{ V}$ pour une alimentation de l'amplificateur opérationnel de $-15\text{ V} / +15\text{ V}$. En fonctionnement linéaire, lorsque le signal d'entrée est sinusoïdal, la tension de sortie doit l'être aussi sauf si $|s|$ devient supérieure aux tensions de saturation pour

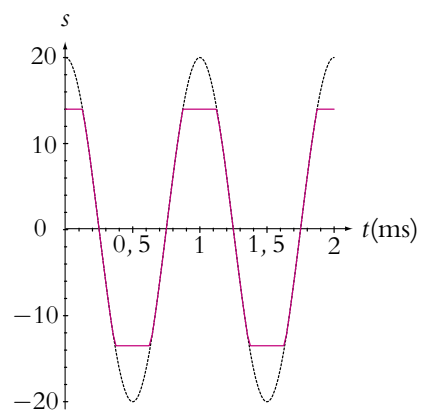


Figure 21.9 Tension attendue en pointillés et tension écrêtée en traits pleins.

certains intervalles de temps. Dans ce cas, le signal sinusoïdal est écrêté comme sur la figure 21.9. La seule solution pour pallier ce problème est de diminuer l'amplitude du signal d'entrée.

3.2 Saturation en courant

Un amplificateur opérationnel ne peut pas délivrer un courant de sortie dont l'intensité $i_{S,AO}$ dépasse une valeur limite $i_{S,AO,max}$. Elle est généralement comprise entre 20 et 30 mA.

On peut mettre en évidence ce phénomène sur tout montage à amplificateur opérationnel fonctionnant en régime linéaire. On se place à la limite de saturation en tension en choisissant des signaux d'entrée tels que le signal de sortie s ait une amplitude proche de la tension de saturation V_{sat} sans la dépasser. On place alors en sortie de l'amplificateur opérationnel une résistance de charge R_C de l'ordre de 10 k Ω : le courant de sortie vaut alors $\frac{V_{sat}}{R_C}$ de l'ordre du milliampère, c'est-à-dire en dessous de la valeur maximale que peut délivrer l'amplificateur opérationnel.

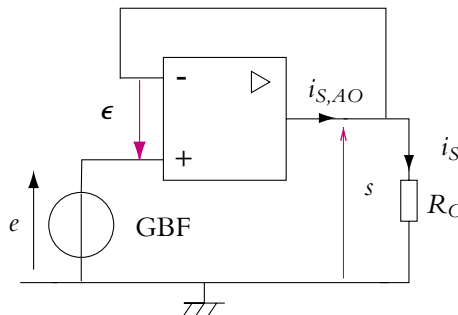


Figure 21.10 Montage pour mettre en évidence la saturation en courant.

Si on diminue la valeur de R_C (en pratique, on utilise un potentiomètre comme résistance R_C), les autres paramètres restant constants, le courant de sortie augmente. À partir d'une certaine valeur de R_C , on observe que le signal de sortie (qui est sinusoïdal) est écrêté comme dans le cas où la tension de sortie dépasse $\pm V_{sat}$. On a donc une saturation. Il ne peut s'agir d'une saturation en tension : on n'a pas fait varier l'amplitude des signaux d'entrée et surtout la valeur du palier de saturation est différente de V_{sat} . Le circuit placé en sortie (ici la résistance R_C) nécessite une intensité plus grande que $i_{S,max}$ et l'amplificateur opérationnel ne peut la délivrer. La tension observée est donc celle aux bornes de R_C : elle est proportionnelle à i_S et est donc écrêtée du fait de la saturation en courant.

Cette mise en évidence peut également permettre de mesurer la valeur de $i_{S,AO,max}$. Dans le cas du montage suiveur, si on suppose l'amplificateur opérationnel idéal, le courant i_- est nul donc $i_s = i_{S,AO}$: en mesurant $i_{S,max}$, on mesure directement

$i_{S,AO,max}$. Il suffit de déterminer la valeur R_{C0} de R_C pour laquelle le phénomène apparaît. Connaissant alors la valeur de la tension maximale (V_{sat}), on en déduit :

$$i_{S,AO,max} = \frac{V_{sat}}{R_{C0}}$$

Attention, si on fait la même mesure avec le montage amplificateur non inverseur, on mesure la somme du courant de sortie de l'amplificateur opérationnel et du courant qui parcourt les résistances R_1 et R_2 . En utilisant les mêmes notations que pour l'exemple précédent, on obtient :

$$i_{s,AO} = i_s - \frac{s}{R_1 + R_2} = i_s - \frac{e}{R_2}$$

avec $i_s = \frac{V_{sat}}{R_{C0}}$ à la limite de saturation en courant.

3.3 Comportement en fréquence de l'amplificateur opérationnel réel

Jusqu'à maintenant, on a considéré que la réponse de l'amplificateur opérationnel ne dépendait pas de la fréquence. En réalité, il en est autrement. Expérimentalement, on constate que, dans la zone de fonctionnement linéaire de l'amplificateur opérationnel, le gain en tension dépend de la fréquence. Ainsi, avec les notations habituelles de la figure 21.3, la relation traduisant le fonctionnement linéaire de l'amplificateur opérationnel s'écrit pour des signaux sinusoïdaux :

$$\underline{s} = \underline{A_d}(\omega) \left(\underline{v_+} - \underline{v_-} \right) \quad \text{avec} \quad \underline{A_d}(\omega) = \frac{A_0}{1 + \frac{j\omega}{\omega_c}} \quad (21.2)$$

A_0 est le gain en tension déjà rencontré qui est supérieur à 10^5 et que l'on prend infini si l'on considère l'amplificateur opérationnel idéal. Sur l'expression de la fonction de transfert ainsi définie, on voit que l'amplificateur opérationnel se comporte comme un filtre passe-bas du premier ordre. Ce comportement est en fait imposé par la structure interne du composant. La pulsation de coupure ω_c est de l'ordre de $100 \text{ rad}\cdot\text{s}^{-1}$.

Ainsi la caractéristique d'un amplificateur opérationnel réel (Cf. figure 21.5) est différente de celle de l'amplificateur opérationnel idéal (Cf. figure 21.6) dans la zone de fonctionnement linéaire. Au lieu d'une droite verticale, la caractéristique est oblique et sa pente $A_d(\omega)$ dépend de la pulsation.

3.4 Stabilité d'un montage à amplificateur opérationnel - Influence du bouclage (PCSI)

L'étude de la stabilité d'un montage à amplificateur opérationnel permet de dire si l'amplificateur opérationnel peut fonctionner en régime linéaire ou s'il est en régime de saturation.

Pour étudier la stabilité d'un montage à amplificateur opérationnel, on fait l'hypothèse que ce dernier fonctionne en régime linéaire et on établit l'équation différentielle liant la tension de sortie s et les tensions d'entrée. Pour cela, on se place le plus souvent dans l'hypothèse de signaux sinusoïdaux et on utilise le passage de la fonction de transfert à l'équation différentielle.

Lors de l'établissement de l'équation différentielle, on ne considère plus l'amplificateur idéal mais on tient compte de son comportement en fonction de la fréquence comme filtre du premier ordre.

À partir de l'équation différentielle, on analyse la stabilité à l'aide des critères donnés au chapitre sur le régime sinusoïdal forcé :

le montage est stable si tous les coefficients de l'équation différentielle sont de même signe.

La conséquence pour un montage à amplificateur opérationnel sera un fonctionnement en régime de saturation si le montage est instable et le fonctionnement en régime linéaire sera possible si le montage est stable.

On verra sur les exemples de montages qui seront traités dans la suite de ce chapitre que l'existence d'un bouclage entre l'entrée inverseuse et la sortie rend possible un fonctionnement de l'amplificateur en régime linéaire.

On admettra les règles suivantes. Si on souhaite s'en convaincre par le calcul, il suffit de faire l'étude décrite ci-dessus pour le montage considéré.

- S'il y a uniquement un bouclage entre l'entrée inverseuse et la sortie, le fonctionnement linéaire est *a priori* possible.
- S'il y a uniquement un bouclage entre l'entrée non inverseuse et la sortie, le fonctionnement linéaire est impossible.
- S'il n'y a aucun bouclage, le fonctionnement linéaire est impossible.

Que faire lorsqu'il y a un bouclage entre l'entrée inverseuse et la sortie et aussi entre l'entrée non inverseuse et la sortie ? Dans ce cas, il faut établir l'équation différentielle et étudier la stabilité du montage. Cependant le fonctionnement est souvent indiqué dans l'énoncé d'un problème ou vérifié expérimentalement.

3.5 Vitesse de balayage ou slew rate

Il existe une limitation de la vitesse de variation de la tension en sortie de l'amplificateur opérationnel. Elle correspond au fait que le temps de réponse de l'amplificateur opérationnel n'est pas négligeable et qu'il y aura un retard dans la réponse.

Cela se traduit par une mauvaise reproduction des variations rapides comme celles indiquées sur la figure 21.11. Ainsi pour un signal carré, le changement brutal de valeur pour la tension ne sera pas immédiat et cela entraîne un phénomène de triangularisation du carré. On parle de *slew rate* ou de *vitesse de balayage limite* en sortie.

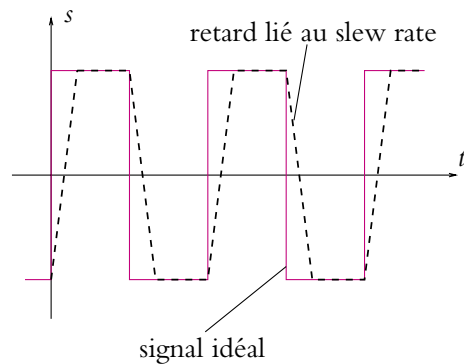


Figure 21.11 Manifestation du slew rate.

Usuellement la vitesse limite est de l'ordre de $0,5 \text{ V} \cdot \mu\text{s}^{-1}$ pour les 741 et de $10 \text{ V} \cdot \mu\text{s}^{-1}$ pour les TL081.

$$\frac{dV_s}{dt} < \left(\frac{dV_s}{dt} \right)_{\max} = \text{SR}$$

Plus la fréquence est élevée, plus le phénomène est visible car on laisse de moins en moins de temps au composant pour réagir.

De plus, la fréquence à laquelle cette limitation intervient dépend de l'amplitude du signal : plus le signal aura une grande amplitude, plus la fréquence maximale d'utilisation (c'est-à-dire en dessous de laquelle aucune déformation du signal n'est visible) sera basse (Cf. paragraphe concernant le montage suiveur).

3.6 Tension de décalage ou offset

On appelle *tension d'offset* ou *tension de décalage* la différence de potentiel qui existe entre les bornes d'entrée inverseuse et non inverseuse de l'amplificateur opérationnel alors qu'on impose un même potentiel à ces deux bornes.

En pratique, si on impose $\epsilon = v_+ - v_- = 0$, on observe une tension de sortie s non nulle. Pour imposer une différence de potentiel nulle entre les deux bornes d'entrée, une solution simple consiste à relier les deux bornes à la masse.

L'existence de cette tension de décalage montre que l'amplificateur opérationnel qu'on supposait symétrique et qui ne devait amplifier que la différence de potentiel imposée à l'entrée ne vérifie pas rigoureusement ces propriétés. Cette différence de potentiel est liée à la réalisation du composant intégré : son origine est l'alimentation de l'amplificateur opérationnel et la dissymétrie du composant.

On modélise cette tension de décalage en intercalant une source idéale de tension à l'entrée de l'amplificateur opérationnel de sorte qu'on ait une tension différentielle

$$\epsilon' = \epsilon + V_d$$

en notant V_d la tension de décalage et ϵ la tension différentielle de l'amplificateur opérationnel idéal. En général, on l'intercale au niveau de la borne d'entrée non inverseuse mais ce choix n'est pas obligatoire.

On a donc la représentation suivante :

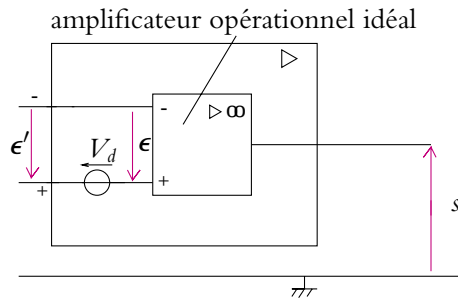


Figure 21.12 Modélisation de la tension d'offset d'un amplificateur opérationnel.

De cette manière, il est tout à fait possible d'imposer $\epsilon' = 0$ et que pour autant $\epsilon = -V_d \neq 0$, ce qui crée une tension de sortie non nulle.

La conséquence sur la caractéristique de transfert est de la décaler parallèlement à l'axe des abscisses :

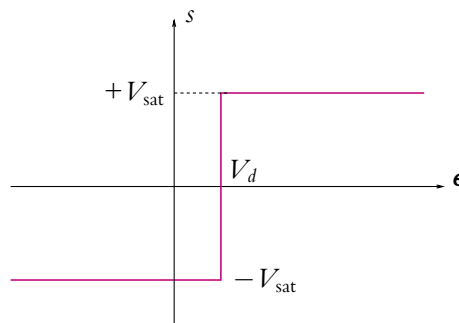


Figure 21.13 Caractéristique de transfert d'un amplificateur opérationnel avec tension d'offset.

On verra en exercice comment il est possible de mesurer cette tension d'offset.

L'inconvénient majeur de la tension d'offset est un risque de saturation de la tension de sortie alors que tous les paramètres du circuit ont été choisis pour que ce type de problème ne se pose pas. On note que la saturation n'est alors pas symétrique : le

phénomène se produit plus tôt pour l'une des saturations que pour l'autre. On risque alors un écrêtage pour la partie haute ou la partie basse du signal suivant le signe de la tension d'offset sans que l'autre soit affectée.

Ce type de problèmes peut se produire lorsqu'on a un fort gain. Il se manifeste quand l'amplitude du signal à amplifier devient du même ordre de grandeur que la tension de décalage V_d . Les deux tensions sont alors amplifiées et la tension de saturation en sortie sera atteinte deux fois plus vite.

En faible gain, la tension d'offset sera souvent négligeable devant la tension d'entrée (car on choisit cette dernière avec des conditions beaucoup moins restrictives). Dans ce cas, le phénomène de saturation en sortie provient essentiellement de la valeur de la tension d'entrée et on n'a pas à se préoccuper de la tension de décalage.

On précisera pour chaque montage particulier qui sera étudié dans la suite du chapitre si la tension d'offset est réellement un inconvénient.

Il existe des méthodes de compensation de cette tension de décalage. En voici un exemple.

On utilise les bornes 1, 4 et 5 du brochage de l'amplificateur opérationnel et on construit un montage potentiométrique afin de « resymétriser » artificiellement le composant. De cette façon, on aura bien une tension de sortie nulle quand on impose une différence de potentiel nulle à l'entrée.

La correction de l'offset s'obtient, par exemple, sur le montage qui a permis sa mesure (Cf. exercice) ou sur un montage où on a simplement relié les deux bornes d'entrée à la masse du circuit. On règle le potentiomètre jusqu'à ce que la tension de sortie soit nulle et ainsi on aura bien corrigé le défaut par rapport au modèle idéal.

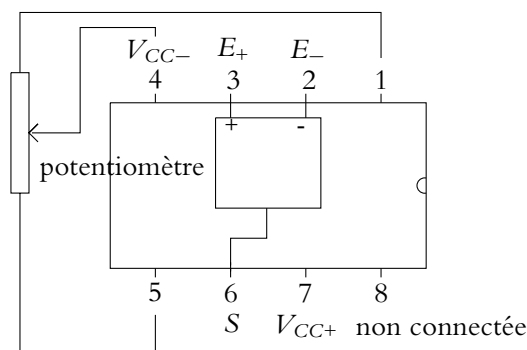


Figure 21.14 Correction de la tension d'offset d'un amplificateur opérationnel.

Sur la figure 21.14, on a noté V_{CC-} et V_{CC+} les tensions -15 V et $+15\text{ V}$ de l'alimentation de l'amplificateur opérationnel.



Attention, ce réglage n'a rien de curatif : il faut laisser le potentiomètre et ne plus toucher à son réglage pour la suite de l'utilisation de l'amplificateur opérationnel. Si on débranche le potentiomètre, on réduit à néant les efforts de réglage de l'offset.

3.7 Courants de polarisation

On appelle *courants de polarisation* les courants circulant dans les bornes d'entrée de l'amplificateur opérationnel.

Dans le modèle de l'amplificateur opérationnel parfait ou idéal, on les suppose nuls mais, en réalité, ils ne le sont pas : on est donc amené à les négliger. Cependant, pour certains montages, il est nécessaire d'en tenir compte.

D'autre part, ils peuvent prendre des valeurs différentes aux entrées inverseuse et non inverseuse en raison de la dissymétrie des composants internes de l'amplificateur opérationnel.

On modélise ces courants à l'aide de sources idéales de courant dont on note I_{0+} et I_{0-} les valeurs des courants de court-circuit. On obtient le schéma :

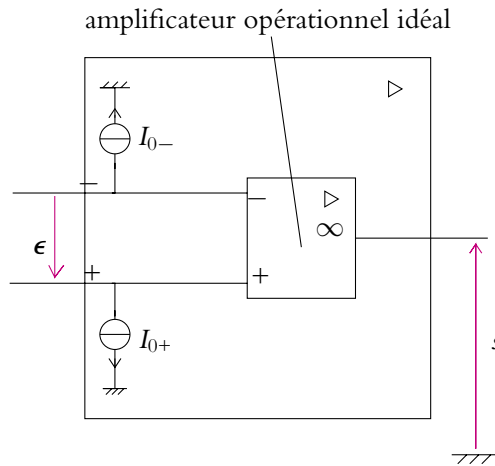


Figure 21.15 Modélisation des courants de polarisation d'un amplificateur opérationnel.

Les conséquences de l'existence de tels courants de polarisation sont également un risque de saturation indésirable. Dans le cas des amplificateurs opérationnels de type TL081, ils sont au maximum de 20 pA et pour les 741, de 500 nA . On aura donc des effets moins notables avec les TL081 qu'avec les 741.

On verra en exercice une méthode pour mesurer les courants de polarisation et comment on peut limiter leurs effets.

3.8 Impédances d'entrée et de sortie

Dans le modèle idéal, l'impédance d'entrée de l'amplificateur opérationnel est infinie et l'impédance de sortie nulle.

En réalité, l'impédance d'entrée est modélisable par une résistance de très grande valeur mais non infinie :

$$R_d = 100 \text{ k}\Omega \text{ à } 10 \text{ G}\Omega$$

De même, l'impédance de sortie est assimilable à une résistance ρ de très faible valeur (de l'ordre de quelques ohms) mais non nulle.

On obtient ainsi le modèle de la figure 21.16 pour l'amplificateur opérationnel en tenant compte des impédances d'entrée et de sortie.

Leur mesure peut s'obtenir par la méthode basée sur le pont diviseur de tension détaillée dans le chapitre sur l'instrumentation. En pratique, ces mesures ne seront pas réalisables du fait des valeurs numériques de R_d et de ρ .

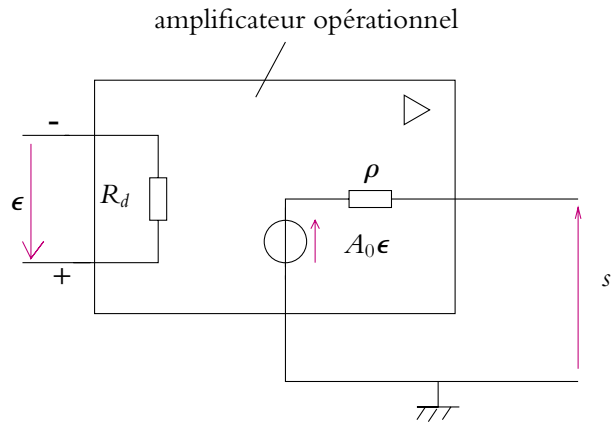


Figure 21.16 Impédances d'entrée et de sortie d'un amplificateur opérationnel.

3.9 Schéma de l'amplificateur opérationnel réel

a) Grandeurs caractéristiques

Caractéristiques	TL081	741
Gain en boucle ouverte	$A_0 = 2.10^5$	$A_0 = 2.10^5$
Produit gain - bande passante	3 MHz	1,5 MHz
Impédance d'entrée différentielle	$Z_d = 1 \text{ G}\Omega$	$Z_d = 1 \text{ M}\Omega$
Tension de décalage d'entrée	$V_{\text{off}} < 15 \text{ mV}$	$V_{\text{off}} < 7,5 \text{ mV}$
Slew rate	$13 \text{ V}.\mu\text{s}^{-1}$	$0,5 \text{ V}.\mu\text{s}^{-1}$
Courants de polarisation	$\simeq 30 \text{ pA}$	$\simeq 80 \text{ nA}$
Courant limite de sortie	20-30 mA	20-30 mA
Puissance maximale	600-700 mW	600-700 mW

b) Modélisation

On n'a représenté que les principaux écarts : tension de décalage et courants de polarisation. On se rappelle que A_d n'est pas infini et surtout qu'il dépend de la fréquence.

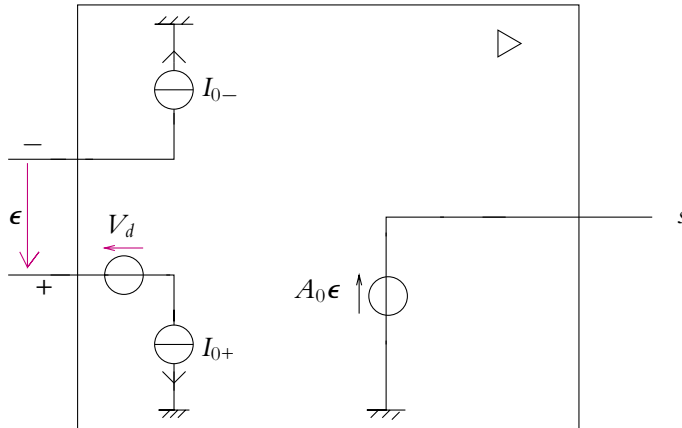


Figure 21.17 Modèle réel d'un amplificateur opérationnel.

4. Montage suiveur

Il s'agit du montage donné par la figure 21.7, déjà présenté au paragraphe 2.1.

4.1 Étude de la stabilité du montage

C'est le plus simple des montages fonctionnant en régime linéaire.

On va établir le résultat en supposant l'amplificateur opérationnel idéal sauf pour la dépendance du gain avec la fréquence.

On suppose les tensions sinusoïdales. On a $\underline{v}_+ = \underline{E}_T$ du fait que les courants de polarisation sont nuls ($i_+ = 0$) et $\underline{v}_- = \underline{s}$. Alors en reportant dans la fonction de transfert de l'amplificateur opérationnel :

$$\underline{A_d} = \frac{\underline{s}}{\underline{v}_+ - \underline{v}_-} = \frac{\underline{s}}{\underline{E}_T - \underline{s}} = \frac{A_0}{1 + j\frac{\omega}{\omega_c}}$$

soit

$$j\frac{\omega}{\omega_c}\underline{s} + (1 + A_0)\underline{s} = A_0\underline{E}_T$$

et l'équation différentielle est :

$$\frac{1}{\omega_c} \frac{ds}{dt} + (1 + A_0)s = A_0 E_T$$

Les coefficients de l'équation différentielle $\frac{1}{\omega_c}$ et $(1 + A_0)$ sont tous les deux positifs : le montage est stable et l'amplificateur opérationnel peut fonctionner en régime linéaire.

4.2 Étude dans l'hypothèse d'un amplificateur opérationnel idéal

L'étude a été effectuée au paragraphe 2.1. On rappelle que, dans ce cas, $s = E_T$.

4.3 Adaptation d'impédance en tension

On peut se demander quelle est l'utilité de ce montage puisque la tension de sortie est égale à la f.e.m. du générateur. Pourquoi n'a-t-on pas branché directement le générateur aux bornes de la résistance de charge R_C ?

Si tel était le cas, on aurait le circuit de la figure 21.18. La tension aux bornes de la charge serait obtenue par un diviseur de tension :

$$s = \frac{R_C}{R_C + \underline{Z}_T} E_T$$

La tension aux bornes de R_C dépendrait de l'impédance interne $\underline{Z}_T$ du générateur.

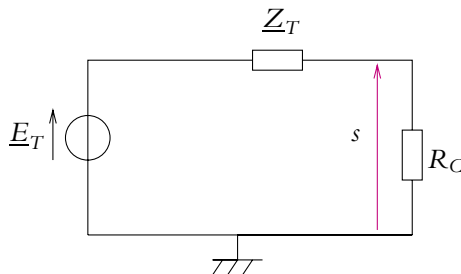


Figure 21.18 Générateur directement relié à la charge.

Pour qu'elle n'en dépende pas, il faudrait qu'elle soit très petite devant la valeur de la charge : $|\underline{Z}_T| \ll R_C$. Or ce n'est pas toujours le cas.

Le montage suiveur permet donc d'avoir une tension aux bornes de la charge indépendante de l'impédance interne du générateur utilisé. Autrement dit, le courant débité par le générateur, égal à i^+ , est quasi-nul. On a ainsi transformé le générateur réel en générateur parfait.

On dit qu'on a réalisé une *adaptation d'impédances pour le transfert de tension*.

4.4 Limitation due à la saturation en tension

On rappelle que la tension de sortie est toujours inférieure en valeur absolue à la tension de saturation V_{sat} . Comme ici $s = E_T$, le montage se comportera effectivement en suiveur sans écrêtage si $|E_T| < V_{\text{sat}}$.

4.5 Impédance d'entrée et impédance de sortie

a) Définition

L'impédance d'entrée est, par définition, égale au rapport entre la tension d'entrée (ici v_+) et l'intensité d'entrée (ici i_+). Dans l'hypothèse où les courants de polarisation sont nuls, $i_+ = 0$ et comme $v_+ = E_T$ a une valeur finie, l'impédance d'entrée est infinie.

Si on tient compte des courants de polarisation, ces derniers sont très faibles et l'impédance d'entrée est donc très grande notamment devant l'impédance Z_T , ce qui permet l'adaptation d'impédance au transfert de tension.

En sortie, le montage peut être représenté par son modèle de Thévenin, la résistance de ce modèle est la résistance de sortie du montage.

On peut alors représenter le montage suiveur par :

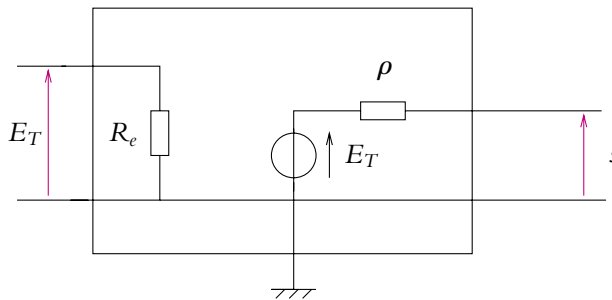


Figure 21.19 Impédances d'entrée et de sortie d'un amplificateur opérationnel.

b) Mesure de l'impédance d'entrée

L'impédance qu'on cherche à déterminer est très grande. Les résistances usuelles n'atteignent pas ces valeurs. On cherche donc simplement à estimer une borne inférieure en s'inspirant de la méthode vue dans le chapitre sur l'instrumentation électrique.

Une difficulté supplémentaire se présente : on ne peut placer un oscilloscope ou un multimètre sur l'entrée du montage car ces appareils ont des impédances d'entrée beaucoup plus faibles que l'amplificateur opérationnel, et c'est l'impédance d'entrée de l'appareil de mesure qui serait mesurée.

On effectue donc une mesure aux bornes du G.B.F. et une mesure de la tension de sortie du montage suiveur, sachant que $s = E_T$.

La méthode consiste à placer une résistance R variable en série avec l'entrée de l'amplificateur opérationnel ou du montage suiveur comme sur le montage 21.20.

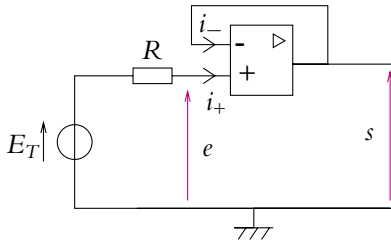


Figure 21.20 Montage pour mesurer l'impédance d'entrée d'un suiveur.

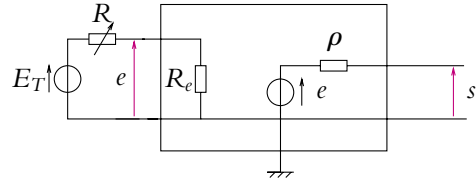


Figure 21.21 Même montage avec le quadripôle équivalent au suiveur.

Sans cette résistance, la mesure de la tension de sortie donnerait :

$$s = e = E_T$$

Avec une résistance R , on aura alors :

$$e = \frac{R_e}{R_e + R} E_T$$

et on mesure en sortie :

$$s = e = \frac{R_e}{R_e + R} E_T$$

En se plaçant à basses fréquences, on augmente R au maximum, par exemple $10^6 \Omega$. On n'observe quasiment aucune variation de s donc de e . Cela signifie que $R_e \gg R$. On en déduit que :

$$R_e > 10^6 \Omega$$

c) Mesure de l'impédance de sortie

Étant donnée la très faible valeur de l'impédance de sortie du montage suiveur, on ne peut effectuer sa mesure mais seulement donner une borne supérieure.

On peut essayer d'appliquer la méthode dite de la demi-tension décrite dans le chapitre sur l'instrumentation électrique.

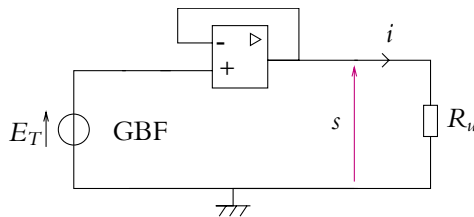


Figure 21.22 Mesure de l'impédance de sortie du montage suiveur.

Quand il n'y a pas de charge R_u , on mesure la tension à vide, à savoir E_T .

On note qu'en réalité, il faut tenir compte de l'impédance d'entrée de l'appareil de mesure. Mais celle-ci (pour un bon appareil de mesure) est très grande et on peut négliger le courant i devant $\frac{s - e}{\rho}$.

Quand on place la résistance R_u , alors on a :

$$s = \frac{R_u}{R_u + \rho} E_T$$

On essaie de faire varier la valeur de R_u pour obtenir une tension dont la valeur soit la moitié de celle obtenue en l'absence de résistance R_u soit

$$\frac{1}{2} = \frac{R_u}{R_u + \rho}$$

ce qui équivaut à

$$R_u = \rho$$

On peut ainsi mettre en évidence que $\rho < 10 \Omega$. Mais il faut faire attention à ce que le courant de sortie ne soit pas saturé (si c'est le cas, il faut diminuer l'amplitude de la tension d'entrée).

4.6 Influence de la tension d'offset

Le gain du suiveur vaut 1 donc $E_T \simeq s$ est de l'ordre du volt. Comme V_d est de l'ordre du millivolt, on aura $E_T \gg V_d$, la tension d'offset sera donc négligeable et le réglage de l'offset sera inutile.

4.7 Influence du slew rate

On suppose maintenant que $E_T = E \sin \omega t$. On a $\frac{dE_T}{dt} = E\omega \cos \omega t$. Pour le montage suiveur $s = E_T$. On en déduit

$$\frac{ds}{dt} = E\omega \cos \omega t$$

Or l'existence d'une vitesse de balayage maximale s'écrit :

$$\left| \frac{ds}{dt} \right| \leq \text{SR}$$

soit

$$E\omega |\cos \omega t| \leq E\omega \leq \text{SR}$$

Ce phénomène n'est pas visible tant que la fréquence du signal vérifie :

$$f \leq \frac{1}{2\pi E} \text{SR}$$

La condition dépend explicitement de l'amplitude E du signal. Il s'agit donc d'une limitation non linéaire.

On constate également que cette condition sur la fréquence sera d'autant plus difficile à réaliser que l'amplitude du signal sera importante. Par exemple, pour une amplitude $E = 10 \text{ V}$, $f \leq \frac{1,3 \cdot 10^5}{2\pi \cdot 10} = 2 \text{ kHz}$ et pour $E = 0,1 \text{ V}$, $f \leq 200 \text{ kHz}$.

D'autre part, cela aura des conséquences sur la forme du signal : les signaux tendent à devenir triangulaires. Cela pose, par exemple, des problèmes lors du tracé des diagrammes de Bode : pour s'affranchir de cette difficulté à une fréquence donnée, il faudra réduire l'amplitude du signal d'entrée. Il est de plus conseillé pour faire l'étude en fréquence de montage à amplificateur opérationnel d'utiliser un TL081 qui a un slew rate beaucoup plus grand que le 741.

On peut mesurer la vitesse maximale de balayage à l'aide de ce montage qu'on alimente par un créneau de forte amplitude sans atteindre la tension de saturation V_{sat} et on mesure la pente de montée ou de descente lors d'un basculement du créneau.

On notera que le basculement d'un créneau correspond aux très grandes fréquences : on est sûr d'être dans le cas où on peut mettre en évidence le phénomène.

5. Montage amplificateur non inverseur

Il s'agit du montage de la figure 21.8 déjà présenté au paragraphe 2.2.

5.1 Étude de la stabilité du montage

On veut établir l'équation différentielle reliant s et e . Pour cela, on suppose les signaux sinusoïdaux pour déterminer la fonction de transfert dans l'hypothèse d'un fonctionnement en régime linéaire de l'amplificateur opérationnel.

$$\underline{s} = \frac{A_0}{1 + \frac{j\omega}{\omega_c}} \left(\underline{v}_+ - \underline{v}_- \right) \quad (21.3)$$

Or ici en remarquant qu'on a un pont diviseur de tension, $\underline{v}_- = \frac{R_1}{R_1 + R_2} \underline{s} = \alpha \underline{s}$ en

posant $\alpha = \frac{R_1}{R_1 + R_2}$. D'autre part, $\underline{v}_+ = \underline{e}$.

En reportant ces expressions dans la relation (21.3), on obtient :

$$\underline{s} = \frac{A_0}{1 + \frac{j\omega}{\omega_c}} (\underline{e} - \alpha \underline{s})$$

et la fonction de transfert :

$$\frac{\underline{s}}{\underline{e}} = \frac{A_0}{1 + \alpha A_0 + j \frac{\omega}{\omega_c}} \quad (21.4)$$

soit

$$(1 + \alpha A_0) \underline{s} + j \frac{\omega}{\omega_c} \underline{s} = \mu_0 \underline{e}$$

On en déduit par l'équivalence $j\omega \Leftrightarrow \frac{d}{dt}$ l'équation différentielle :

$$(1 + \alpha A_0) s + \frac{1}{\omega_c} \frac{ds}{dt} = A_0 e$$

Les coefficients de l'équation différentielle homogène associée sont tous positifs donc tous de même signe : le système est stable et l'amplificateur opérationnel peut fonctionner en régime linéaire.

5.2 Influence du comportement en fréquence de l'amplificateur opérationnel

a) Fonction de transfert

On l'obtient à partir de la relation (21.4) donnant la fonction de transfert du montage. On peut réécrire cette relation sous la forme canonique d'un filtre passe-bas du premier ordre en faisant apparaître un terme $1 + j \frac{\omega}{\omega'_c}$ au dénominateur soit :

$$\underline{H} = \frac{\underline{s}}{\underline{e}} = \frac{\frac{A_0}{1 + \alpha A_0}}{1 + j \frac{\omega}{(1 + \alpha A_0) \omega_c}}$$

On retrouve bien la forme canonique d'un filtre passe-bas du premier ordre :

$$\underline{H} = \frac{A'_0}{1 + j \frac{\omega}{\omega'_c}}$$

avec

$$A'_0 = \frac{A_0}{1 + \alpha A_0} \quad \text{et} \quad \omega'_c = \omega_c (1 + \alpha A_0)$$

Le montage amplificateur non inverseur se comporte donc comme un filtre passe-bas du premier ordre lorsqu'on tient compte du comportement fréquentielle de l'amplificateur opérationnel.

b) Le produit gain - bande passante est constant

On remarque que :

$$A'_0 \omega'_c = A_0 \omega_c$$

A_0 et A'_0 sont les gains statiques respectivement de l'amplificateur opérationnel et du montage amplificateur non inverseur tandis que ω_c et ω'_c sont les fréquences de coupure des filtres passe-bas correspondants.

On dit que le produit gain-bande passante est conservé.



Attention : ce résultat est vrai pour ce montage particulier de l'amplificateur non inverseur mais ce n'est pas le cas pour tous les montages à amplificateur opérationnel.

On peut remarquer en lisant le tableau des caractéristiques des amplificateurs opérationnels fourni au paragraphe a) que les fabricants donnent le produit gain-bande passante du composant. Il est alors possible de déterminer la fréquence de coupure des deux amplificateurs opérationnels usuels en divisant le produit gain - bande passante obtenu par A_0 :

- pour le TL081 : $f_c = \frac{3 \cdot 10^6}{2 \cdot 10^5} = 15 \text{ Hz}$,
- pour le 741 : $f_c = \frac{1,5 \cdot 10^6}{2 \cdot 10^5} = 7,5 \text{ Hz}$.

c) Diagramme de Bode en gain

On peut s'intéresser aux diagrammes de Bode en gain et notamment à leur allure asymptotique :

$$G_{AO,dB} = 20 \log \mu \quad \text{et} \quad G_{NI,dB} = 20 \log H$$

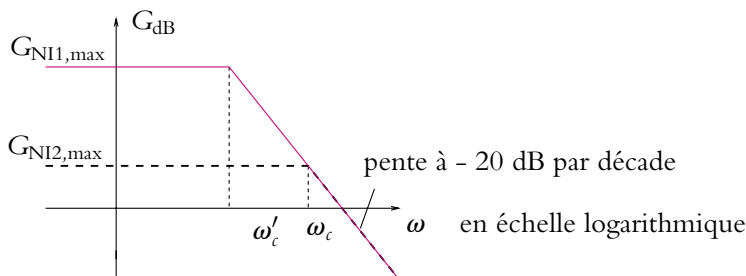


Figure 21.23 Diagramme asymptotique en gain de l'amplificateur non inverseur.

Pour l'amplificateur opérationnel, l'asymptote en haute fréquence est :

$$20 \log A_0 - 20 \log \frac{\omega}{\omega_c}$$

et pour l'amplificateur non inverseur, c'est $20 \log A'_0 - 20 \log \frac{\omega}{\omega'_c}$. Si on remplace A'_0 et ω'_c par leurs expressions en fonction de ω_0 et A_0 , on trouve que les asymptotes sont les mêmes. C'est une des conséquences du fait que le produit gain - bande passante est constant.

En basses fréquences, les asymptotes sont respectivement : $20 \log A_0$ et $20 \log A'_0$.

On obtient donc le diagramme de Bode asymptotique de la figure 21.23.

On peut vérifier que lorsqu'on fait tendre A_0 vers l'infini, on retrouve le cas de l'amplificateur opérationnel idéal soit $\underline{H} = 1 + \frac{R_2}{R_1}$.

5.3 Validité de l'approximation de l'amplificateur opérationnel idéal

Peut-on alors encore utiliser l'hypothèse de l'amplificateur opérationnel idéal ? Pour qu'elle soit valable, il faut que la fréquence du signal sinusoïdal soit dans la bande passante du montage. Par exemple, pour l'amplificateur non inverseur, si on choisit comme valeurs de résistances $R_1 = 1 \text{ k}\Omega$ et $R_2 = 10 \text{ k}\Omega$ alors $\alpha = \frac{1}{11}$ et :

$$A'_0 = \frac{A_0}{1 + \alpha A_0} = \frac{2 \cdot 10^5}{1 + \frac{2 \cdot 10^5}{11}} \simeq 11$$

On retrouve donc que le transfert statique du montage est égal à la fonction de transfert trouvée pour ce montage dans l'hypothèse de l'amplificateur opérationnel idéal :

$$A'_0 = 11 = 1 + \frac{R_2}{R_1}$$

Avec ces valeurs de résistances et pour le TL081, la fréquence de coupure est $f'_c = f_c \frac{A_0}{A'_0} = 15 \frac{2 \cdot 10^5}{11} \simeq 270 \text{ kHz}$.

Jusqu'à une fréquence d'un dixième de la fréquence de coupure, le gain est le même qu'en continu donc il ne faut prendre en compte les écarts à l'idéalité de l'amplificateur opérationnel qu'aux hautes fréquences. Les études qui ont été faites en considérant l'amplificateur opérationnel idéal restent donc valables sur une large gamme de fréquences.

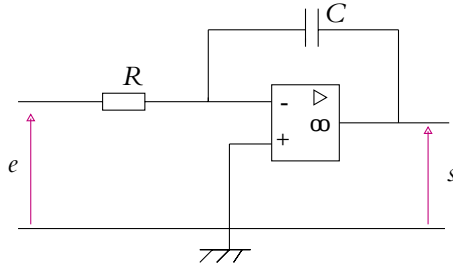
5.4 Influence de la tension d'offset

Dans le cas où le gain est important de l'ordre de 10^3 , la tension d'entrée doit être inférieure à $\frac{V_{sat}}{1000} \simeq 15 \text{ mV}$. En se plaçant en limite de fonctionnement par rapport à la saturation en tension, la tension d'entrée e est alors du même ordre de grandeur que la tension d'offset V_d . On aura donc intérêt à régler la tension d'offset à 0 pour éviter une saturation en tension.

6. Montages intégrateur et pseudo-intégrateur

6.1 Montage intégrateur théorique et amplificateur opérationnel idéal

Le montage théorique d'un intégrateur à amplificateur opérationnel est le suivant :



On suppose que l'amplificateur opérationnel est idéal et qu'il fonctionne en régime linéaire. On néglige l'influence de la fréquence sur le comportement de l'amplificateur opérationnel. On peut donc écrire : $\underline{\epsilon} = 0$. Or $\underline{v}_+ = 0$. En écrivant la loi des nœuds en termes de potentiel à l'entrée inverseuse de l'amplificateur opérationnel (ceci est possible du fait que les courants de polarisation sont nuls), on obtient :

$$\frac{\underline{e} - \underline{v}_-}{R} + jC\omega (\underline{s} - \underline{v}_-) = 0$$

La condition $\underline{\epsilon} = \underline{v}_+ - \underline{v}_- = 0$, donne $\underline{v}_- = 0$ puisque $\underline{v}_+ = 0$, d'où :

$$\underline{e} + jRC\omega \underline{s} = 0$$

ce qui se traduit en termes d'équation différentielle par :

$$e(t) + RC \frac{ds}{dt} = 0$$

ou encore

$$s(t) = -\frac{1}{RC} \int_{t_0}^t e(t) dt + s(t_0)$$

Le signal de sortie est donc l'intégration du signal d'entrée à un coefficient RC près.

En pratique, ce montage ne permet pas de réaliser l'intégration du signal d'entrée : on observe un phénomène de saturation.

► **Remarque** : le passage en notation complexe n'est pas indispensable pour établir la relation entre $e(t)$ et $s(t)$. En effet, la loi des nœuds en termes de potentiel s'écrit, en régime quelconque :

$$\frac{e - v_-}{R} = C \frac{d(v_- - s)}{dt}$$

Compte tenu du fait que $v_- = 0$, il vient :

$$e = -RC \frac{ds}{dt} \quad \text{soit} \quad s(t) = -\frac{1}{RC} \int_{t_0}^t e(t) dt + s(t_0)$$

6.2 Influence de la tension d'offset

Lors de la présentation de la tension d'offset, on a dit que la compensation de cet écart par rapport au modèle idéal était important lors d'un fort gain. Or ici le gain dépend de la fréquence et vaut :

$$G = \frac{1}{RC\omega}$$

Si on admet qu'il faut un gain de l'ordre de 10^3 pour qu'il faille tenir compte de l'offset, on aura :

$$\frac{1}{RC\omega} \leq 10^3 \quad \text{soit} \quad \omega \geq \frac{1}{10^3 RC}$$

Si on choisit $R = 10 \text{ k}\Omega$ et $C = 1 \text{ }\mu\text{F}$, on devra vérifier : $\omega \geq \frac{1}{10^3 \times 10^4 \times 10^{-6}} = 0,1$ ou une fréquence $f \geq 1,6 \cdot 10^{-2} \text{ Hz}$. Il n'y aura donc pas ici à se préoccuper de la tension d'offset à cause du gain.

En revanche, un autre problème se pose : celui de la dérive. En effet, il y aura toujours la tension d'offset qui est continue et sera donc en permanence intégrée. L'intégration du signal d'entrée sera donc « décalée » puis on ne l'observera plus car il y aura saturation du fait que l'intégration de la tension d'offset aura atteint la tension de saturation.

Par le calcul, on a :

$$\begin{cases} v_+ = V_d \\ v_- = \frac{jC\omega \underline{e} + \frac{\underline{e}}{R}}{jC\omega + \frac{1}{R}} = \frac{jRC\omega \underline{e} + \underline{e}}{1 + jRC\omega} \end{cases}$$

En régime linéaire, on a : $\underline{e} = v_+ - v_- = 0$ soit

$$V_d = \frac{jRC\omega \underline{e} + \underline{e}}{1 + jRC\omega}$$

et

$$jRC\omega \underline{e} = (1 + jRC\omega) V_d - \underline{e}$$

En notation réelle, on obtient :

$$s = -\frac{1}{RC} \int_{t_0}^t e(t) dt + \frac{V_d}{RC} (t - t_0) + s(t_0)$$

C'est le terme $\frac{V_d}{RC} t$ qui impose une dérive et sature l'amplificateur opérationnel.

Par exemple, si on suppose que $\forall t, e(t) = 0$ alors

$$s = \frac{V_d}{RC}t + s(0)$$

Ainsi pour une tension de sortie initialement nulle, le montage arrivera à saturation au bout d'un temps :

$$t_1 = RC \frac{V_{\text{sat}}}{V_d}$$

Si on choisit $R = 10 \text{ k}\Omega$ et $C = 10 \text{ nF}$ et si on suppose que la tension de décalage est $V_d \simeq 10 \text{ mV}$, la saturation sera atteinte pour $t_1 \simeq 0,15 \text{ s}$. En conclusion, la tension d'offset provoque une rapide saturation du montage intégrateur.

Il sera indispensable ici d'effectuer le réglage de l'offset ou de modifier le montage pour s'en affranchir en utilisant un montage pseudo-intégrateur (Cf. paragraphe 6.4).

6.3 Influence des courants de polarisation

On étudie maintenant l'effet du courant de polarisation $i_- = I_0$ en supposant que la tension d'offset a été réglée à 0 et que $i_+ = 0$. En supposant toujours que $\forall t, e(t) = 0$, la loi des noeuds à l'entrée inverseuse s'écrit alors :

$$\frac{-v_-}{R} = i_- + C \frac{d(v_- - s)}{dt}$$

Or ici $v_- = v_+ = 0$ et $i_- = I_0$ d'où :

$$\frac{ds}{dt} = \frac{I_0}{C} \Rightarrow s = \frac{I_0}{C} t + s(0)$$

Ainsi pour une tension de sortie initialement nulle, le montage arrivera à saturation au bout d'un temps :

$$t_2 = C \frac{V_{\text{sat}}}{I_0}$$

Avec les mêmes valeurs de R et C qu'au paragraphe précédent, on a les temps suivants pour atteindre la saturation :

- pour le 741 pour lequel $I_0 = 100 \text{ nA}$: $t_2 = \frac{15 \cdot 10^{-8}}{100 \cdot 10^{-9}} = 1,5 \text{ s}$,
- pour le TL081 pour lequel $I_0 = 30 \text{ pA}$: $t_2 = \frac{15 \cdot 10^{-8}}{3 \cdot 10^{-12}} = 5 \cdot 10^4 \text{ s}$.

En conclusion, les courants ne provoquent une saturation rapide que pour le 741. On pourra ainsi utiliser le montage intégrateur théorique avec un TL081 simplement en réglant la tension de décalage à 0. Cependant ce réglage n'est jamais parfait et la dérive de la tension de sortie finit toujours par apparaître. Une solution pour pallier le problème de dérive est d'utiliser le montage pseudo-intégrateur (Cf. paragraphe 6.4).

6.4 Montage pseudo-intégrateur

Le montage intégrateur proposé précédemment présente une saturation en sortie due aux écarts par rapport au modèle idéal de l'amplificateur opérationnel.

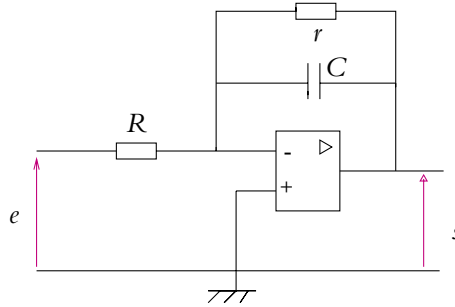


Figure 21.24 Montage pseudo-intégrateur.

Le montage de la figure 20.24 permet de s'affranchir de ce phénomène de saturation. On suppose que l'amplificateur opérationnel est idéal et qu'il fonctionne en régime linéaire. On peut donc utiliser la relation $\epsilon = v_+ - v_- = 0$ soit $v_+ = v_- = 0$. En négligeant les courants de polarisation, on écrit la loi des nœuds en termes de potentiel à l'entrée inverseuse en notant $\underline{Z}$ l'impédance équivalente à l'association en parallèle de C et r :

$$\frac{e}{R} + \frac{s}{\underline{Z}} = 0$$

Or

$$\underline{Z} = \frac{1}{\frac{1}{r} + jC\omega} = \frac{r}{1 + jrC\omega}$$

on en déduit donc :

$$s = \frac{r}{R(1 + jrC\omega)} e$$

On obtient donc un filtre passe-bas du premier ordre qui n'aura un comportement intégrateur que pour des pulsations telles que $\omega \gg \frac{1}{rC}$.

Par conséquent, les constantes (qui correspondent à un signal de fréquence nulle) comme la tension de décalage et les courants de polarisation ne sont pas intégrées par ce montage. On s'affranchit ainsi du phénomène de dérive. En contrepartie, on ne peut pas intégrer les signaux de n'importe quelle fréquence : seuls ceux dont les fréquences vérifient $f \gg \frac{1}{2\pi rC}$ pourront être intégrés. C'est pourquoi on parle de *pseudo-intégrateur*.

De plus, si on choisit la fréquence de coupure très faible (rC de l'ordre de quelques hertz), le montage ne laisse passer que la composante continue du signal d'entrée c'est-à-dire sa valeur moyenne : un filtre passe-bas de fréquence de coupure très faible est donc un opérateur « valeur moyenne ».

7. Complément : autres exemples de montages simples

7.1 Amplificateur inverseur

Le montage étudié est celui de la figure 21.25. On cherche à exprimer s en fonction de e .

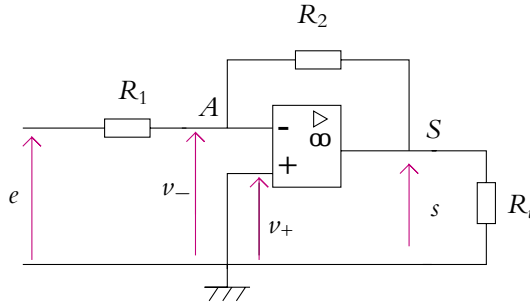


Figure 21.25 Montage inverseur

On fait les hypothèses suivantes :

- l'amplificateur opérationnel est idéal :
 - $i_+ = i_- = 0$ A (1),
 - gain infini (2),
- on vérifie qu'il y a un bouclage entre v_- et s donc un fonctionnement en régime linéaire est *a priori* possible (3).

$$(2) + (3) \Rightarrow v_+ = v_-$$

La loi des nœuds exprimée en termes de potentiel en A s'écrit, compte tenu de $i_- = 0$:

$$\frac{e - v_-}{R_1} = \frac{v_- - s}{R_2} \quad (21.5)$$

Sachant que $v_+ = 0$ et $v_+ = v_-$, l'équation (21.5) donne :

$$\frac{s}{e} = -\frac{R_2}{R_1} \quad (21.6)$$

Le montage porte le nom d'*amplificateur inverseur* : inverseur en raison du signe « $-$ » entre s et e et amplificateur en raison du rapport R_2/R_1 (on note que si $R_2 < R_1$, la tension d'entrée n'est pas amplifiée mais atténuée).

► **Remarques** La résistance d'utilisation R_u n'intervient pas, aussi ne sera-t-elle plus représentée. Il ne faut cependant pas qu'elle soit choisie en dessous d'une certaine valeur pour éviter la saturation en courant.

Cela ne sert à rien d'appliquer la loi des nœuds à la sortie s car on introduit une inconnue supplémentaire, le courant de sortie de l'amplificateur opérationnel, sur lequel on n'a aucune information.

La relation (21.6) est valable tant que $|s| < V_{\text{sat}}$.

7.2 Montage sommateur

Le montage étudié est celui de la figure 21.26. On cherche à exprimer s en fonction de e_1 et e_2 .

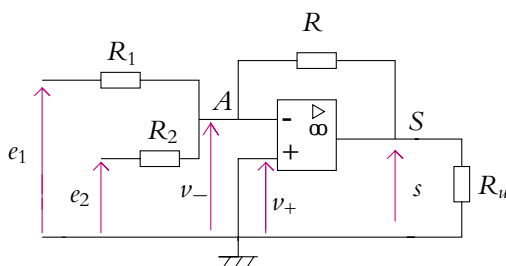


Figure 21.26 Montage sommateur

On fait les hypothèses suivantes :

- l'amplificateur opérationnel est idéal :
 - $i_+ = i_- = 0$ A (1),
 - gain infini (2),
- on vérifie qu'il y a un bouclage entre v_- et s donc un fonctionnement en régime linéaire est *a priori* possible (3).

$$(2) + (3) \Rightarrow v_+ = v_-$$

La loi des nœuds exprimée en termes de potentiel en A s'écrit, compte tenu de $i_- = 0$:

$$\frac{e_1 - v_-}{R_1} + \frac{e_2 - v_-}{R_2} = \frac{v_- - s}{R} \quad (21.7)$$

Sachant que $v_+ = 0$ et $v_+ = v_-$, l'équation (21.7) donne :

$$s = -R \left(\frac{e_1}{R_1} + \frac{e_2}{R_2} \right) \quad (21.8)$$

V_s est donc égale à la somme des tensions e_1 et e_2 pondérée par les résistances du circuit. La relation (21.8) est valable tant que $|s| < V_{\text{sat}}$.

Il peut y avoir plus de deux résistors connectés à l'entrée inverseuse. Dans ce cas, l'expression (21.8) se généralise par :

$$s = -R \sum_{j=1}^n \frac{e_n}{R_n} \quad (21.9)$$

en notant e_n la tension appliquée au résistor de résistance R_n .

7.3 Montage dérivateur

Le montage étudié est celui de la figure 21.27. On cherche à exprimer s en fonction de e .

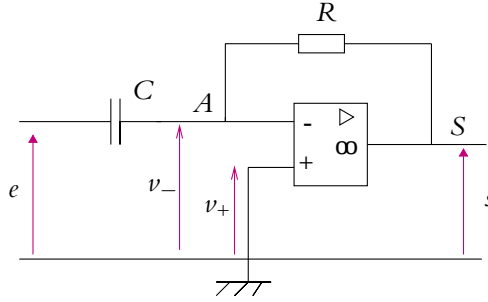


Figure 21.27 Montage dérivateur

On fait les hypothèses suivantes :

- l'amplificateur opérationnel est idéal :
 - $i_+ = i_- = 0$ A (1),
 - gain infini (2),
- on vérifie qu'il y a un bouclage entre v_- et s donc un fonctionnement en régime linéaire est *a priori* possible (3).

$$(2) + (3) \Rightarrow v_+ = v_-$$

La loi des nœuds exprimée en termes de potentiel en A s'écrit, compte tenu de $i_- = 0$:

$$C \frac{d(e - v_-)}{dt} = \frac{v_- - s}{R} \quad (21.10)$$

Avec $v_+ = 0$ et $v_- = v_+$, l'équation (21.10) donne :

$$s = -RC \frac{de}{dt} \quad (21.11)$$

La relation (21.11) est valable tant que $|s| < V_{\text{sat}}$.

On peut également effectuer le calcul directement avec les impédances : la loi des nœuds exprimée en termes de potentiel en A s'écrit :

$$jC\omega \left(\underline{e} - \underline{v_-} \right) = \frac{\underline{v_-} - \underline{s}}{R} \quad (21.12)$$

soit, puisque $\underline{v_+} = \underline{v_-} = 0$:

$$\underline{s} = -jRC\omega \underline{e} \quad (21.13)$$

On remarque la cohérence des deux expressions démontrées puisque la multiplication par $j\omega$ correspond à une dérivation.

8. Compareurs à amplificateur opérationnel

Les *compareurs* sont des circuits fonctionnant en régime non linéaire. La tension de sortie ($\pm V_{\text{sat}}$) dépend du résultat de la comparaison de la tension d'entrée avec une tension de référence, d'où le nom donné à ces montages.

8.1 Comparateur simple

C'est le plus simple des montages compareurs à amplificateur opérationnel : il est représenté sur la figure 21.28. On impose une tension E sur l'entrée non inverseuse : $v_+ = E$. De plus, on a à l'entrée inverseuse : $v_- = e$.

Puisqu'il n'y a aucun bouclage, le montage ne peut fonctionner qu'en régime non linéaire et on a :

- si $e > E$ alors $v_- > v_+$ et $s = -V_{\text{sat}}$,
- si $e < E$ alors $v_- < v_+$ et $s = V_{\text{sat}}$.

On en déduit la caractéristique de transfert du montage donnant s en fonction de e , qui est représentée sur la figure 21.29.

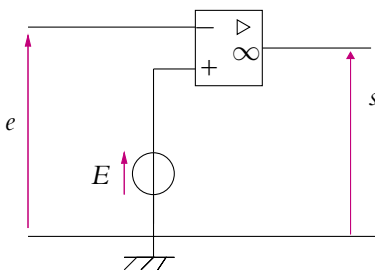


Figure 21.28 Comparateur simple inverseur.

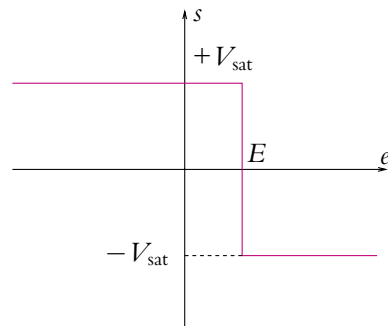


Figure 21.29 Caractéristique de transfert d'un comparateur simple inverseur.

La tension de sortie est positive quand la tension d'entrée est inférieure à la valeur seuil et inversement, la sortie est donc « inversée » par rapport à l'entrée. Cela justifie le terme de comparateur inverseur donné à ce montage.

Il existe des variantes de ce montage : en particulier, on peut imposer $v_- = E$ et envoyer e sur l'entrée non inverseuse. Dans ce cas, la caractéristique est inversée par rapport à celle de la figure 21.29. La tension de sortie est alors négative quand la tension d'entrée est inférieure à la valeur seuil et inversement, la sortie n'est donc pas « inversée » par rapport à l'entrée. Cela justifie le terme de comparateur non inverseur donné à ce montage.

Une autre variante consiste à relier l'une des entrées à la masse au lieu de lui imposer une tension E : l'entrée est alors comparée à 0.

Ce montage est très simple mais peu performant. **En particulier, il est très sensible au bruit c'est-à-dire aux parasites.** Par exemple, si ce système commande l'ouverture d'une vanne de remplissage d'une grande cuve d'eau lorsque la tension s est positive et sa fermeture lorsque la tension s est négative. Un capteur mesure la hauteur d'eau h et délivre une tension proportionnelle à h : $e = kh$ en notant k une constante. Soit h_0 la hauteur telle que $h_0 = \frac{E}{k}$. Lorsque $h > h_0$ alors $e > E$ et $s < 0$ d'après l'étude précédente donc la vanne se ferme. Au contraire, si $h < h_0$, la vanne s'ouvre. Tout semble donc fonctionner parfaitement. Mais si le vent, par exemple, crée des vaguelettes à la surface du liquide, le niveau de l'eau peut passer successivement au-dessous et en dessus de h_0 . La vanne va donc s'ouvrir puis se refermer très rapidement puis s'ouvrir de nouveau... Le résultat sera un vieillissement prématuré du système.

Ainsi les parasites sur le signal d'entrée ont une influence importante sur le signal de sortie de ce montage, ce qui n'est pas le cas du montage qu'on va étudier au paragraphe suivant.

On peut d'ailleurs visualiser ce phénomène en travaux pratiques en envoyant en entrée, par exemple, un signal sinusoïdal de 50 Hz auquel on ajoute un signal sinusoïdal de 2 kHz d'amplitude plus faible. On choisit dans ce cas $E = 0$ V.

8.2 Comparateur à hystérésis

Le montage étudié est le suivant :

a) Étude de la stabilité du montage

Avant d'analyser le comportement du montage de la figure 21.30, on va étudier sa stabilité et montrer que le bouclage sur l'entrée non inverseuse rend ce montage instable.

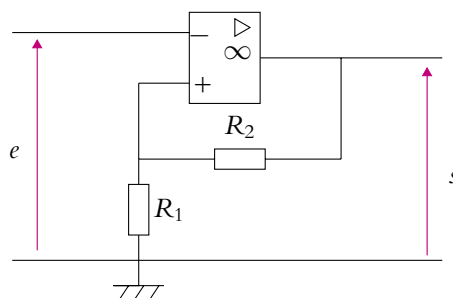


Figure 21.30 Comparateur à hystérésis.



Attention : ce montage ressemble à celui de l'amplificateur non inverseur mais ici le bouclage se fait entre l'entrée non inverseuse et la sortie.

On utilise toujours la même méthode : on suppose que les signaux sont sinusoïdaux et que l'amplificateur opérationnel fonctionne en régime linéaire pour déterminer l'équation différentielle reliant s et e . Pour cela, on cherche la fonction de transfert du montage.

En régime linéaire :

$$\underline{s} = \frac{A_0}{1 + \frac{j\omega}{\omega_c}} (\underline{v}_+ - \underline{v}_-) \quad (21.14)$$

Or ici on obtient à l'aide d'un pont diviseur de tension :

$$\underline{v}_+ = \frac{R_1}{R_1 + R_2} \underline{s} = \alpha \underline{s} \quad \text{avec} \quad \alpha = \frac{R_1}{R_1 + R_2}$$

D'autre part, on a : $\underline{v}_- = \underline{e}$. En reportant dans la relation (21.14), on obtient :

$$\underline{s} = \frac{A_0}{1 + \frac{j\omega}{\omega_c}} (\alpha \underline{s} - \underline{e})$$

On en déduit la fonction de transfert :

$$\frac{\underline{s}}{\underline{e}} = - \frac{A_0}{1 - \alpha A_0 + j \frac{\omega}{\omega_c}}$$

et l'équation différentielle par l'équivalence : $j\omega \Leftrightarrow \frac{d}{dt}$:

$$\frac{1}{\omega_c} \frac{dv_s}{dt} + (1 - \alpha A_0) s = -A_0 e$$

Sachant que pour des montages courants $\alpha > 10^{-5}$ (par exemple, pour $R_1 = 1 \text{ k}\Omega$ et $R_2 = 10 \text{ k}\Omega$: $\alpha = 1/11$), le terme $1 - \alpha A_0$ est négatif puisque $A_0 > 2.10^5$. Les coefficients de l'équation différentielle homogène associée sont de signes opposés donc le système est instable. Le signal de sortie saturera rapidement vers $\pm V_{\text{sat}}$ en fonction des conditions initiales.

La sortie n'est plus proportionnelle à l'entrée, le montage ne fonctionne plus en régime linéaire ni en régime sinusoïdal forcé mais en régime de saturation.

b) Analyse du fonctionnement

D'après le paragraphe précédent, l'amplificateur opérationnel fonctionne en régime non linéaire : la tension de sortie est $\pm V_{\text{sat}}$.

Le raisonnement pour ce type de fonctionnement est différent de celui utilisé lorsque l'amplificateur opérationnel fonctionne en régime linéaire. Dans le cas d'un fonctionnement en régime non linéaire, on connaît la tension de sortie puisqu'elle est égale à $\pm V_{\text{sat}}$. Contrairement aux montages linéaires, on part donc de la sortie pour déterminer l'état de l'entrée.

Premier cas

On suppose tout d'abord que $s = V_{\text{sat}}$.

Ceci est possible tant que $v_+ > v_-$. Or ici en négligeant les courants de polarisation ($i_- = 0$), on obtient la tension v_+ à l'aide des résultats sur les ponts diviseurs de tension : $v_+ = \frac{R_1}{R_1 + R_2} V_{\text{sat}}$. D'autre part, $v_- = e$.

Cet état est donc possible tant que $e < \frac{R_1}{R_1 + R_2} V_{\text{sat}}$ ou $e < V_B$ en posant $V_B = \frac{R_1}{R_1 + R_2} V_{\text{sat}}$.

Second cas

On suppose maintenant que $s = -V_{\text{sat}}$.

Ceci est possible tant que $v_+ < v_-$. La tension v_+ vaut maintenant :

$$v_+ = -\frac{R_1}{R_1 + R_2} V_{\text{sat}}$$

et on a toujours $v_- = e$.

Cet état est donc possible tant que $e > -\frac{R_1}{R_1 + R_2} V_{\text{sat}}$ ou $e > -V_B$ avec la même définition de V_B que précédemment.

Caractéristique de transfert

Pour déterminer la caractéristique de transfert donnant s en fonction de e , on suppose que le signal d'entrée est sinusoïdal d'amplitude supérieure à V_B .

D'après l'étude précédente, on s'aperçoit que si $-V_B < e < V_B$, les deux états $\pm V_{\text{sat}}$ de la tension de sortie sont possibles. Il faut donc choisir le signal e tel qu'à l'instant initial, seul un des états soit possible en sortie. On choisit, par exemple, le signal représenté sur la figure 21.31.

En suivant l'évolution de e , on construit simultanément le signal de sortie s (Cf. figure 21.32) ainsi que la caractéristique de transfert donnant s en fonction de e représentée sur la figure 21.33. Les points servant à la construction portent les mêmes numéros sur les trois figures. Les étapes sont les suivantes :

1. À $t = 0$, $e < -V_B < V_B$ donc seul l'état $s = +V_{\text{sat}}$ est possible (point 1).
2. Le signal d'entrée e augmente et atteint la valeur $-V_B$ (point 2). Les deux états de sortie deviennent *a priori* possibles d'après l'étude du paragraphe précédent. Mais comme $v_+ - v_- = V_B - e = 2V_B > 0$, la condition pour avoir $s = +V_{\text{sat}}$ est toujours vérifiée : il n'y a pas basculement et s reste à $+V_{\text{sat}}$.
3. Le signal d'entrée e augmente encore et atteint la valeur V_B (point 3). Seul l'état $s = -V_{\text{sat}}$ devient possible : la tension de sortie bascule de $+V_{\text{sat}}$ à $-V_{\text{sat}}$.
4. Le signal d'entrée e continue d'augmenter jusqu'à son maximum (point 4), la tension de sortie ne change pas puisque la comparaison entre les potentiels des entrées inverseuse et non inverseuse de l'amplificateur donne toujours la même conclusion.
5. La tension d'entrée e diminue et atteint V_B (point 5). De nouveau, les deux valeurs de la tension de sortie deviennent *a priori* possibles. Comme au point 2, on détermine : $v_+ - v_- = -V_B - e = -2V_B < 0$ donc la tension de sortie s reste à $-V_{\text{sat}}$.
6. La tension d'entrée e continue de diminuer jusqu'à atteindre $-V_B$ (point 6) alors seul l'état $s = +V_{\text{sat}}$ redevient possible et le système bascule en sortie de la tension $-V_{\text{sat}}$ à $+V_{\text{sat}}$.
7. Enfin la tension d'entrée e diminue jusqu'à son minimum et le système revient à la même situation qu'au point 1.

On obtient une caractéristique pour laquelle le chemin est différent pour chaque basculement. Ce type de caractéristique porte le nom de *cycle d'hystérésis* du grec *υστερειν* (husterein) qui signifie « être en retard ». On parle aussi de *trigger ou déclencheur de Schmidt*.

Le montage étudié ici est le comparateur à hystérésis inverseur pour les mêmes raisons que pour le comparateur simple. Le comparateur à hystérésis non inverseur existe également.

On notera qu'il est fondamental d'indiquer par des flèches le sens dans lequel le cycle est décrit. Seules les commutatrices verticales ne peuvent être décrites que dans un sens, les parties horizontales peuvent l'être dans les deux.

Ce type de comparateur est plus performant que le comparateur simple. En effet, comme les tensions de basculement sont différentes, il est moins sensible au bruit.

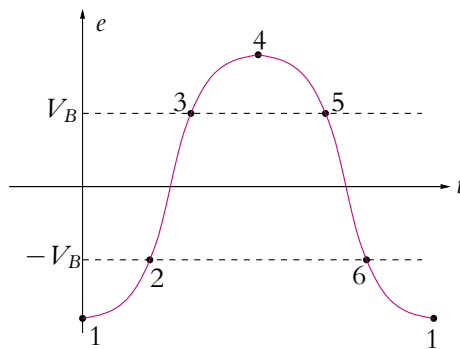


Figure 21.31 Signal d'entrée du comparateur à hystérésis.

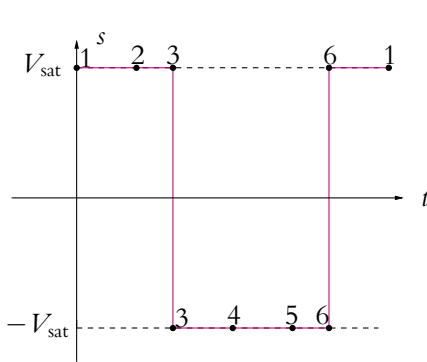


Figure 21.32 Chronogramme de la tension de sortie.

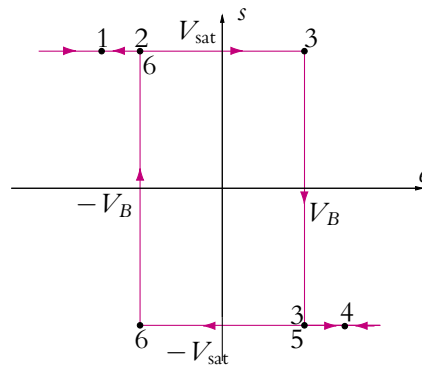


Figure 21.33 Caractéristique de transfert du comparateur à hystérésis inverseur.

D'autre part, les circuits fonctionnant sur ce principe d'hystérésis **peuvent être utilisés comme mémoire électronique**. En effet, une mémoire est un circuit qui garde une information lorsque la tension d'entrée est nulle, c'est-à-dire si on n'applique plus de signal à l'entrée. Or il existe toujours des parasites si bien que la tension d'entrée ne sera pas tout à fait nulle. Le comparateur simple basculera suivant le signe du parasite alors que le comparateur à hystérésis restera dans l'état dans lequel il se trouve sauf si bien sûr l'amplitude du parasite devient supérieure à V_B en valeur absolue, ce qui est rare.

9. Complément : multivibrateur astable

L'une des applications importante du comparateur à hystérésis est la réalisation de système oscillant délivrant des signaux créneaux de période facilement réglable. C'est le principe de base d'un générateur basse fréquence qui fournit des signaux en créneaux qu'on intègre en triangle puis qu'on transforme éventuellement en sinusoïde.

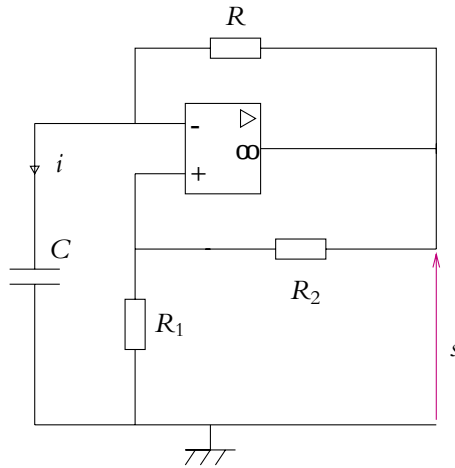


Figure 21.34 Multivibrateur astable

On ne donne pas dans la suite un circuit utilisé dans les G.B.F. mais on étudie le principe d'un oscillateur sur un circuit simple appelé *multivibrateur astable à amplificateur opérationnel*.

Le montage du circuit étudié est représenté sur la figure 21.34. On remarque qu'il y a deux bouclages entre les entrées et les sorties. On constate expérimentalement que ce montage ne fonctionne qu'en mode saturé. On reconnaît la présence du bouclage du comparateur à hystérésis étudié au paragraphe précédent.

9.1 Analyse du fonctionnement

On suppose tout d'abord que la sortie est en saturation haute, $s = +V_{\text{sat}}$, ce qui signifie que la tension v_- est inférieure à la tension v_+ . Dans ce cas, le condensateur de capacité C et la résistance R en série sont soumis à la tension positive V_{sat} , ce qui va entraîner une charge du condensateur à travers la résistance. La tension v_- va augmenter et devenir supérieure à v_+ provoquant le basculement de la tension de sortie de $+V_{\text{sat}}$ à $-V_{\text{sat}}$. Alors le condensateur va se décharger à travers la résistance et v_- va diminuer jusqu'à devenir inférieure à v_+ . On va de nouveau assister à un basculement à la sortie. Et ainsi de suite...

9.2 Mise en équation du phénomène

On suppose qu'à $t = 0$, la sortie vient de basculer de $-V_{\text{sat}}$ à $+V_{\text{sat}}$ (point 2 de la figure 21.33) c'est-à-dire d'après l'étude précédente que $v_- = -V_B$. Dans ce cas, puisqu'on néglige les courants de polarisation ($i_+ = 0$), on applique un diviseur de tension entre l'entrée non inverseuse et la sortie :

$$v_+ = \frac{R_1}{R_1 + R_2} V_{\text{sat}} = V_B$$

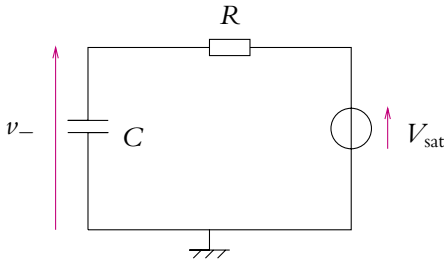


Figure 21.35 Montage équivalent lorsque $s = V_{\text{sat}}$.

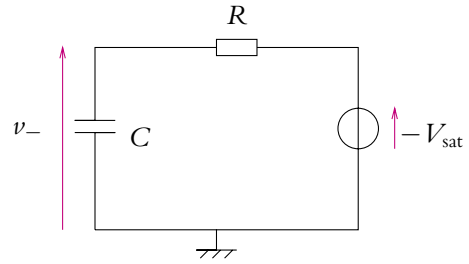


Figure 21.36 Montage équivalent lorsque $s = -V_{\text{sat}}$.

Puisque les courants de polarisation sont négligés, tout se passe comme si la résistance et le condensateur étaient en série avec un générateur de tension parfait de f.e.m. V_{sat} (Cf. figure 21.35). L'équation différentielle de ce circuit est :

$$v_- + RC \frac{dv_-}{dt} = V_{\text{sat}}$$

de solution générale :

$$v_- = A \exp\left(-\frac{t}{\tau}\right) + V_{\text{sat}}$$

où $\tau = RC$ et A est une constante qu'on détermine avec la condition initiale $v_-(t=0) = -V_B$. On trouve alors l'expression de v_- :

$$v_- = -(V_B + V_{\text{sat}}) \exp\left(-\frac{t}{\tau}\right) + V_{\text{sat}} \quad (21.15)$$

Cette solution est valable tant que $v_- < V_B$ soit pour $0 < t < t_1$ avec t_1 tel que :

$$v_-(t_1) = V_B \Rightarrow t_1 = \tau \ln \frac{V_{\text{sat}} + V_B}{V_{\text{sat}} - V_B} = \tau \ln \frac{2R_1 + R_2}{R_2}$$

À l'instant t_1 , la tension de sortie bascule de $+V_{\text{sat}}$ à $-V_{\text{sat}}$ (point 3 sur la figure 21.33).

Le raisonnement est le même que précédemment sauf que le circuit R, C est maintenant soumis à la tension $-V_{\text{sat}}$ (Cf. figure 21.36) et que $v_+ = -V_B$. L'équation différentielle est alors :

$$v_- + \tau \frac{dv_-}{dt} = -V_{\text{sat}}$$

de solution générale :

$$v_- = A' \exp\left(-\frac{t}{\tau}\right) - V_{\text{sat}}$$

La constante A' est cette fois-ci déterminée à l'instant t_1 pour lequel $v_- = V_B$. On en déduit l'expression de v_- :

$$v_- = (V_B + V_{\text{sat}}) \exp\left(-\frac{t - t_1}{\tau}\right) - V_{\text{sat}} \quad (21.16)$$

Cette expression est valable jusqu'à l'instant t_2 pour lequel la tension v_- devient inférieure à $v_+ = -V_B$. Le système bascule de nouveau en sortie à V_{sat} (point 6 sur la figure 21.33).

Le raisonnement est alors à reprendre au début. Puisque le système est revenu à son point de départ, la période T du phénomène est égale à t_2 que l'on détermine en écrivant que $v_-(t_2) = -V_B$ soit :

$$T = t_2 = 2\tau \ln \frac{V_{\text{sat}} + V_B}{V_{\text{sat}} - V_B} = 2t_1 \quad (21.17)$$

L'allure des signaux obtenus sur l'entrée inverseuse et sur la sortie est représentée sur la figure 21.37.

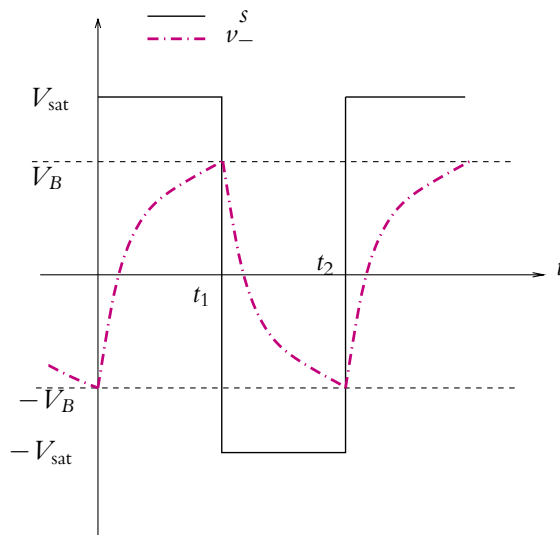


Figure 21.37 Chronogrammes du multivibrateur astable.

On a donc fabriqué un générateur de créneau de période proportionnelle à RC donc facilement réglable avec une résistance variable. L'énergie du système provient de l'alimentation continue de l'amplificateur opérationnel qui n'apparaît pas sur le schéma électrique.

9.3 Principe de fonctionnement d'un générateur basse fréquence

Un générateur basse fréquence délivre des signaux créneaux, triangulaires et sinusoïdaux.

Comme on l'a signalé au paragraphe précédent, il suffit de fabriquer un générateur de créneaux puis d'intégrer le signal créneau par un circuit intégrateur pour obtenir un signal triangulaire. Un circuit de ce type est proposé en exercice.

Pour obtenir un signal sinusoïdal, on pourrait utiliser en sortie un filtre passe-bande pour filtrer le signal triangulaire. Cependant, il faudrait accorder la fréquence de résonance du filtre à chaque changement de fréquence de l'oscillateur. Cette méthode est donc très efficace pour un montage de fréquence fixe mais pas dans le cas d'un générateur de fréquence variable. On utilise alors un circuit appelé *conformateur à diodes*. Son but est « d'arrondir » le signal triangulaire pour le transformer en signal ressemblant à une sinusoïde.

A. Applications directes du cours

1. Saturations

On visualise à l'oscilloscope les signaux représentés sur la figure ci-dessous, ces signaux sont obtenus en sortie d'un amplificateur opérationnel.

1. Proposer des montages pouvant conduire à de telles observations.
2. Pour chaque solution, quels sont les "défauts" de l'amplificateur opérationnel mis en évidence ? Justifier.

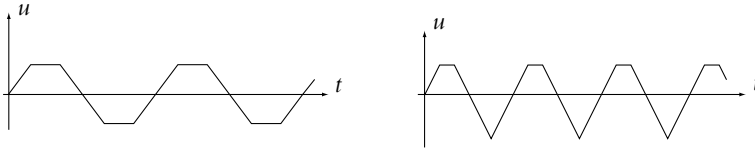


Figure 21.38

2. Mesure de la tension de décalage

On peut mesurer la tension d'offset grâce au montage ci-contre.

1. Montrer que, si on néglige la tension d'offset, la tension s est nulle.
2. On modélise la tension d'offset comme dans le cours. Établir l'expression de s dans ces conditions.
3. En déduire que la mesure de s permet de mesurer la tension d'offset.

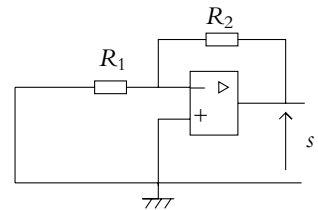


Figure 21.39

3. Courants de polarisation

1. Pour mesurer les courants de polarisation, on peut utiliser le montage suivant :

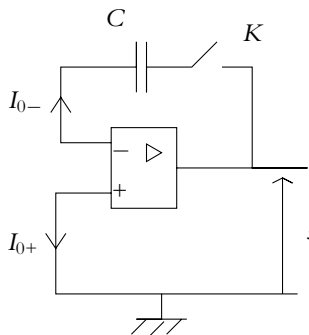


Figure 21.40

On suppose que la capacité C est initialement déchargée. Montrer que la mesure du temps nécessaire pour que le signal de sortie de l'amplificateur opérationnel atteigne une tension U donnée permet de déterminer la valeur des courants de polarisation. On suppose que les courants de polarisation vérifient $I_{0+} = I_{0-}$.

2. Proposer une méthode simple pour compenser les courants de polarisation.
3. Une autre solution consiste à ce que les deux entrées de l'amplificateur opérationnel voient les mêmes impédances quand elles sont reliées à la masse. Comment doit-on choisir par exemple la valeur de la résistance R du montage suivant pour qu'il en soit ainsi ?

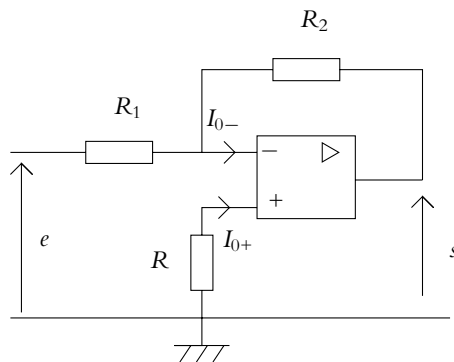


Figure 21.41

4. Impédance d'entrée

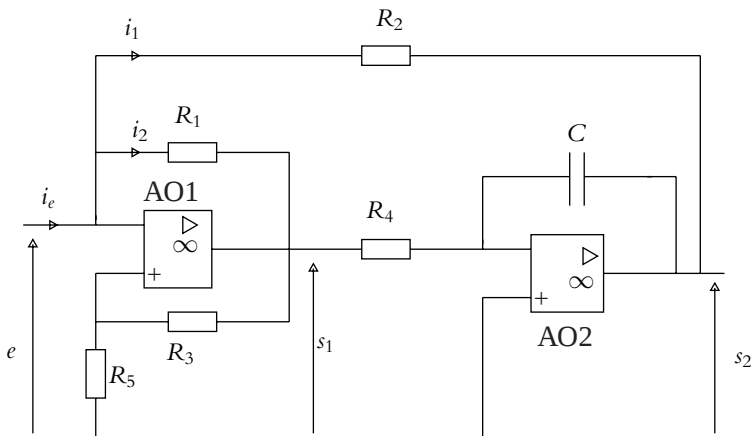


Figure 21.42

On suppose que les amplificateurs opérationnels sont idéaux.
Déterminer l'impédance d'entrée du montage.

5. Exemple de filtre de Butterworth

Soit le montage ci-contre.

On suppose que l'amplificateur opérationnel est idéal.

1. Calculer la fonction de transfert de ce montage.
2. Déterminer la condition que doivent vérifier C_1 et C_2 pour que le module de la fonction de transfert s'écrive :

$$|H| = \frac{1}{\sqrt{1 + \left(\frac{\omega}{\omega_0}\right)^4}}$$

Donner alors la valeur de ω_0 en fonction de R et de C_1 .

3. Étudier les variations du gain. En déduire le tracé réel et asymptotique du gain en décibels en fonction de $\log \omega$.
4. Définir et calculer la bande passante du filtre.
5. Étudier et tracer le déphasage en fonction de $\log \omega$.

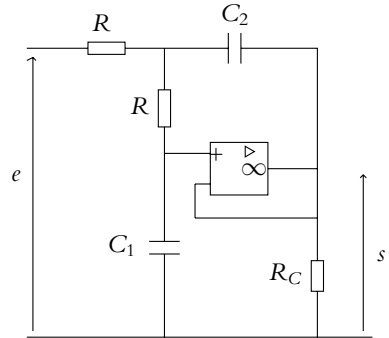


Figure 21.43

6. Amplificateur et puissance, d'après ICNA 2005

On considère le montage suivant où les amplificateurs opérationnels sont supposés idéaux et fonctionner en régime linéaire.

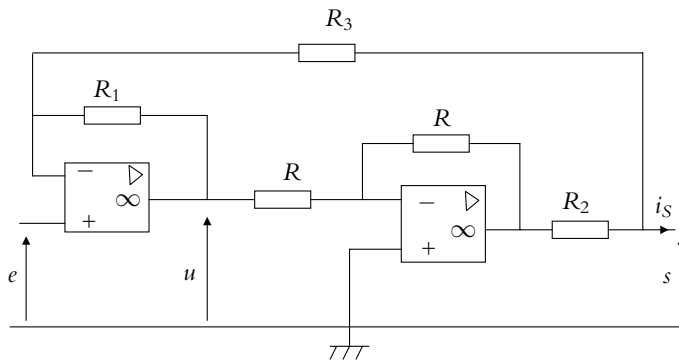


Figure 21.44

1. Rappeler les hypothèses de l'amplificateur opérationnel idéal.
2. Définir le régime linéaire de l'amplificateur opérationnel.
3. Établir l'expression du courant de sortie en court-circuit qu'on notera $I_{s,CC}$.
4. Déterminer l'expression de la tension de sortie s en circuit ouvert.
5. En déduire que le montage vu de la sortie est équivalent à un générateur de Norton dont on précisera les caractéristiques.

6. Soit le montage ci-contre pour lequel on donne $R_4 = 1,0 \text{ k}\Omega$, $R_5 = 10 \text{ k}\Omega$ et $C = 0,10 \text{ }\mu\text{F}$.

Établir l'expression de l'impédance complexe du circuit notée $\underline{Z}$.

7. En déduire sa norme et donner sa valeur numérique lorsque la tension v est sinusoïdale de pulsation $\omega = 1000 \text{ rad.s}^{-1}$ et de valeur efficace 100 V.

8. Établir l'expression de la tangente du déphasage φ introduit par cette impédance.

9. En déduire l'expression du déphasage φ .

10. Donner l'expression de la puissance consommée dans cette impédance ainsi que sa valeur numérique.

11. Mêmes questions pour la puissance consommée dans R_4 .

12. Mêmes questions pour la puissance consommée dans C .

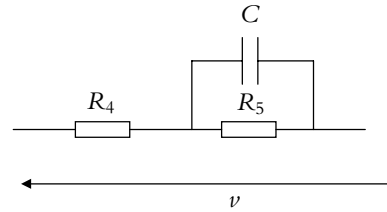


Figure 21.45

B. Exercices et problèmes

1. Résonance avec un montage électronique, d'après Centrale TSI 2003

- Rappeler la caractéristique statique d'un amplificateur opérationnel.
- Donner les hypothèses que doit vérifier un amplificateur opérationnel pour être idéal.
- Qu'appelle-t-on régime linéaire de l'amplificateur opérationnel ?
- Déduire de ce qui précède la relation particulière pour le régime linéaire d'un amplificateur opérationnel idéal.
- Soit le montage suivant :

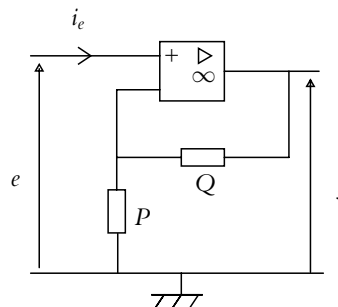


Figure 21.46

Établir la fonction de transfert de ce montage en supposant que l'amplificateur opérationnel est idéal et fonctionne en régime linéaire.

6. Que vaut le courant d'entrée i_e ?

7. Établir que le montage est équivalent au schéma suivant :

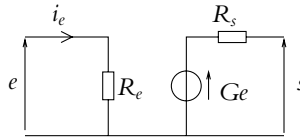


Figure 21.47

On donnera les expressions de R_e , R_s et G .

8. On insère ce type de montage dans le circuit suivant :

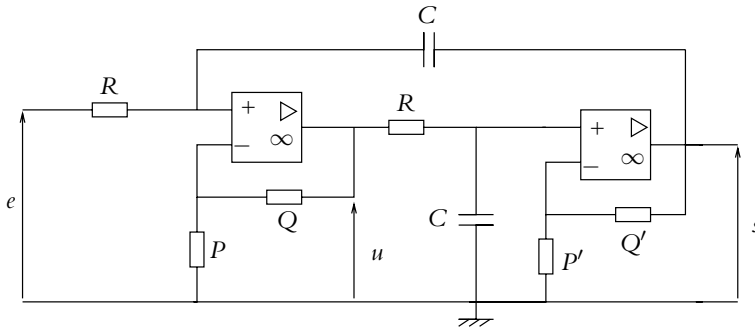


Figure 21.48

Si on veut un gain G de 1 et G' proche de 2, proposer un jeu raisonnable de valeurs pour les résistances P , Q , P' et Q' . Dans la suite, sauf indication contraire, on fera l'étude générale sans tenir compte des valeurs qui viennent d'être proposées pour G et G' .

9. Déterminer la fonction de transfert du montage en fonction de G , G' , R , C et ω .
10. En déduire la valeur du gain en tension en fonction de $x = \frac{\omega}{\omega_0}$ avec $\omega_0 = \frac{1}{RC}$.
11. Étudier les variations du gain en tension en fonction de la fréquence. On distinguera plusieurs cas en fonction des valeurs du produit GG' et on donnera les valeurs maximales du gain en tension en fonction du produit GG' .
12. Donner l'expression approchée du maximum du gain en tension lorsque G vaut 1 et G' est proche de 2 en lui étant inférieur.
13. Calculer dans le cas général la bande passante du montage.
14. Donner son expression approchée lorsque G vaut 1 et G' est proche de 2 en lui étant inférieur et en déduire l'expression de l'acuité du montage en fonction du produit GG' .
15. Donner les allures possibles du gain en tension en fonction de $x = \frac{\omega}{\omega_0}$.
16. Expliquer comment on peut facilement faire varier la valeur de l'acuité de ce montage.

17. On suppose $R = 100 \text{ k}\Omega$, $C = 1,00 \text{ }\mu\text{F}$, $Q' = 8,00 \text{ k}\Omega$, $P' = 10,0 \text{ k}\Omega$, $Q = 1,00 \text{ k}\Omega$, $P = 100 \text{ k}\Omega$. Calculer les valeurs de ω_0 , de G , de G' , de GG' , du gain maximal et du rapport x_0 de la pulsation de résonance à ω_0 .

18. On envoie un signal rectangulaire périodique de période $T = 1,00 \text{ ms}$ prenant la valeur E_0 avec $E_0 = 1,00 \text{ V}$ lorsque $0 \leq t \leq aT$ et 0 lorsque $aT \leq t \leq T$ avec $a = 0,20$. Justifier l'allure du signal de sortie dont on donnera toutes les valeurs caractéristiques (littéralement et numériquement).

2. Réalisation d'un filtre, d'après ESIGETEL MP 2000

1. On considère tout d'abord le circuit suivant composé d'une résistance R et d'un condensateur de capacité C alimentant une charge assimilable à une résistance R' . On note e la tension d'entrée et s celle de sortie.

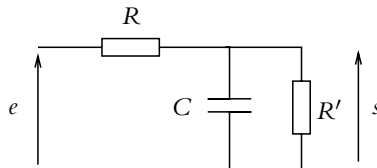


Figure 21.49

Dans un premier temps, la tension d'entrée est une tension continue E .

- Établir l'équation différentielle vérifiée par s .
 - Sachant que pour $t < 0$, le circuit a atteint un régime continu et qu'on commence l'alimentation du circuit par la tension continue E à $t = 0$, que peut-on dire de la valeur de la tension s à cet instant ?
 - En déduire la résolution de l'équation différentielle en s .
 - Représenter la tension $s(t)$ en la justifiant brièvement.
 - Exprimer la tension u aux bornes de R .
 - En déduire la représentation de $u(t)$.
2. On remplace la tension continue par une tension sinusoïdale $e(t) = E \cos \omega t$. On s'intéresse au régime permanent.
- Calculer la fonction de transfert du circuit.
 - Quelle est la nature du filtre obtenu ?
 - L'écrire sous forme canonique et donner les expressions du gain H_0 à fréquence nulle et de la fréquence de coupure f_0 .
 - Tracer, en le justifiant, le diagramme de Bode en gain. On donnera les tracés réel et asymptotique.
 - Même question pour le diagramme de Bode en phase.
3. On cherche à améliorer ce circuit en utilisant un amplificateur opérationnel.

- Rappeler les hypothèses d'un amplificateur opérationnel idéal.
- Qu'appelle-t-on régime linéaire de l'amplificateur opérationnel ?
- Quelle est la conséquence du caractère idéal sur le régime linéaire de l'amplificateur opérationnel ? Dans la suite, on se placera dans ces conditions.
- Déterminer la fonction de transfert du montage suivant :

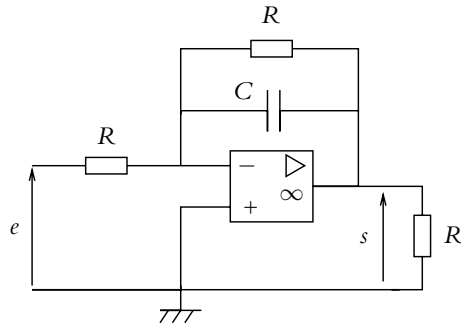


Figure 21.50

- Comparer le diagramme de Bode de ce filtre avec le précédent.
 - En déduire l'intérêt de ce nouveau montage.
4. On remplace alors la tension sinusoïdale par une tension créneau de période T avec $e(t) = E$ pour $0 \leq t \leq \frac{T}{2}$ et $e(t) = -E$ pour $-\frac{T}{2} \leq t \leq 0$.
- Calculer la décomposition en séries de Fourier de ce signal. On rappelle que celle-ci s'écrit

$$\frac{a_0}{2} + \sum_{n=1}^{+\infty} (a_n \cos(n\omega t) + b_n \sin(n\omega t))$$

avec $a_n = \frac{2}{T} \int_{-\frac{T}{2}}^{+\frac{T}{2}} e(t) \cos(n\omega t) dt$ et $b_n = \frac{2}{T} \int_{-\frac{T}{2}}^{+\frac{T}{2}} e(t) \sin(n\omega t) dt$.

- Donner, en la justifiant, l'expression générale de $s(t)$.
- Le montage précédent peut-il se comporter en dérivateur ?
- Même question pour un comportement intégrateur.
- En déduire la condition portant sur le choix de R et C pour obtenir une sortie quasi-triangulaire ?

3. Wattmètre électronique, d'après Banque PT 2006

On se propose d'étudier le fonctionnement d'un wattmètre électronique comprenant un capteur de tension, un capteur de courant, un multiplieur et un filtre passe-bas. Aucune connaissance préalable du wattmètre n'est nécessaire.

1. Soit un dipôle d'impédance $\underline{Z} = R + jX$ alimenté par une tension sinusoïdale $v(t) = V_M \cos(\omega t)$ et parcouru par un courant $i(t) = I_M \cos(\omega t - \varphi)$. On utilise la convention récepteur.

Donner les expressions de l'amplitude I_M et du déphasage φ en fonction de V_M , R et X .

Quel est le signe de φ si le dipôle est inductif?

2. Calculer la puissance moyenne P reçue par le dipôle en fonction de V_M , I_M et φ .

3. Pour une impédance $\underline{Z} = 5 + 10j$ exprimée en ohms et une amplitude $V_M = 320$ V, donner les valeurs numériques de I_M , φ et P .

4. Pour mesurer la puissance P absorbée par la charge, on utilise un wattmètre dont le schéma fonctionnel est le suivant : la tension $v(t)$ est envoyée dans un capteur de tension qui délivre en sortie une tension $v_a(t) = k_a v(t)$; l'intensité est envoyée dans un capteur de courant délivrant en sortie une tension $v_b(t) = k_b i(t)$. Ces deux signaux $v_a(t)$ et $v_b(t)$ entrent dans un multiplieur produisant la tension $v_1(t) = K v_a(t) v_b(t)$. La tension $v_1(t)$ passe ensuite dans un moyennage permettant d'obtenir une tension "lissée" $v_2(t)$ telle que la valeur moyenne de v_2 est l'opposée de celle de v_1 .

Établir l'expression de la valeur moyenne de $v_2(t)$ et montrer qu'elle est proportionnelle à la puissance P . On explicitera la constante de proportionnalité k_m .

5. On mesure une valeur moyenne de $-1,5$ V pour v_2 . Sachant que $k_a = 0,02$, $k_b = 10,0$ mV.A⁻¹ et $K = 2,00$ V⁻¹, donner la valeur de la puissance P mesurée.

6. Pour obtenir la valeur moyenne du signal $v_1(t)$, on utilise un filtre passe-bas dont le montage dit de Rauch est donné ci-dessous. Établir la fonction de transfert $\underline{H}$ de ce montage.

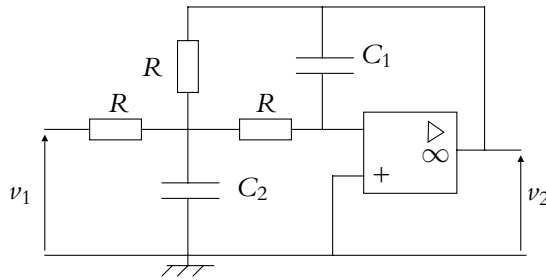


Figure 21.51

7. L'écrire sous la forme canonique $\underline{H} = \frac{H_0}{1 + 2mj \frac{\omega}{\omega_0} - \left(\frac{\omega}{\omega_0}\right)^2}$. Expliciter m , H_0 et ω_0 en

fonction des paramètres du circuits.

8. On souhaite obtenir une fréquence $f_0 = 5,0$ Hz avec un coefficient $m = \frac{\sqrt{2}}{2}$ et $R = 470$ k Ω . Calculer les valeurs des capacités C_1 et C_2 .

9. Établir le sens de variation du gain.

10. Préciser le diagramme asymptotique du diagramme de Bode en gain.

11. Déterminer l'expression du déphasage φ .
 12. Établir son sens de variation.
 13. Donner les valeurs limites du déphasage.
 14. En déduire le tracé du diagramme de Bode.
 15. Donner, sans nouveau calcul, l'équation différentielle reliant v_1 et v_2 en utilisant les constantes ω_0 et m .
 16. On applique à l'entrée du filtre passe-bas la sortie du multiplieur $v_1(t)$ sachant qu'en entrée du wattmètre on envoie $v(t) = V_M \cos(\omega t)$ et $i(t) = I_M \cos(\omega t - \varphi)$ avec $V_M = 320$ V, $f = 50$ Hz, $I_M = 20$ A et $\varphi = 1,0$ rad. Montrer qu'en régime forcé, la tension du sortie du filtre passe-bas est l'opposé de la valeur moyenne de $v_1(t)$ à laquelle on superpose une composante alternative. Déterminer l'amplitude A_M de cette composante alternative et donner sa valeur numérique.
- En déduire la précision relative de la mesure de la puissance P .
17. Le capteur de tension est un transformateur dont le primaire comprend N_1 spires et le secondaire N_2 spires. On suppose le transformateur parfait soit $\frac{u_2}{u_1} = \frac{i_1}{i_2} = \frac{N_2}{N_1}$ en indiquant par 1 les valeurs du primaire et par 2 celles du secondaires, le primaire et le secondaire étant en convention récepteur.
- Déterminer le rapport $\frac{N_2}{N_1}$ pour obtenir une tension efficace de 5,0 V au secondaire pour une tension efficace de 230 V au primaire.
18. En déduire le nombre de spires du secondaire sachant que $N_1 = 500$.
 19. Pour prélever la tension aux bornes de dipôle d'impédance $\underline{Z}$, le primaire est branché en parallèle sur ce dipôle. La résistance d'entrée du multiplieur valant $R_e = 10,0$ k Ω , calculer la puissance moyenne P_e consommée par le capteur de tension avec les valeurs numériques précédentes.

4.

Oscillateur à pont de Wien, d'après ENSTIM 2006

On considère un filtre de Wien constitué d'un pont diviseur de tension dont les deux impédances sont d'une part l'association en série d'une résistance et d'une capacité et d'autre part de l'association en parallèle des mêmes éléments. La tension de sortie est prélevée aux bornes de l'association en parallèle.

1. On alimente le filtre par une tension sinusoïdale $e = E \cos \omega t$. Établir la fonction de transfert du filtre et la mettre sous la forme $\underline{H} = \frac{A}{1 + jQ \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)}$. On donnera les valeurs de A , Q et ω_0 en fonction de R et C .
2. Déterminer les expressions asymptotiques du gain.
3. Exprimer le déphasage et donner ses valeurs asymptotiques.
4. En déduire le diagramme de Bode de ce filtre.

5. Établir l'équation différentielle vérifiée par le système lorsqu'on l'alimente par une tension quelconque.
6. On couple ce filtre à un amplificateur opérationnel parfait. On supposera que ce dernier fonctionne en régime linéaire. Qu'appelle-t-on amplificateur opérationnel parfait ou idéal? A quoi correspond l'hypothèse de fonctionnement linéaire de l'amplificateur opérationnel?

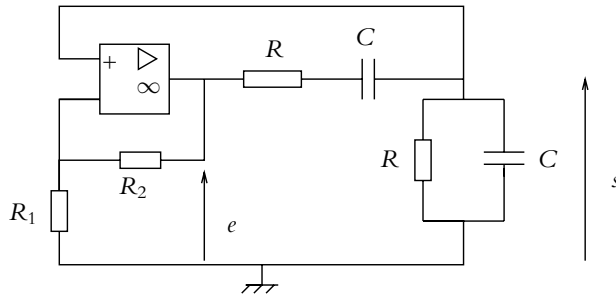


Figure 21.52

7. Établir la relation entre le potentiel à l'entrée inverseuse de l'amplificateur opérationnel et la tension e en fonction de R_1 et R_2 .
8. En déduire l'équation différentielle vérifiée par s .
9. À quelle condition la solution de cette équation est-elle purement sinusoïdale?
10. Que se passe-t-il si $1 - Q \left(1 + \frac{R_2}{R_1} \right) < 0$? On précisera notamment le mode de fonctionnement de l'amplificateur opérationnel dans ces conditions.

5. Aspect électrique de la technique du *patch clamp*, d'après Agro-Véto 2005

Inventé par E. Neher et B. Sakmann, le *patch clamp* ou électrophysiologie cellulaire est la technique de référence pour l'étude des canaux ioniques. Chaque canal est traversé par un flux ionique élevé (de l'ordre de 10^6 ions par second) et génère un courant électrique. Cette technique consiste à enregistrer l'activité électrique de la membrane cellulaire.

Soit un milieu homogène conducteur de conductivité électrique σ . Une charge q de masse m se déplace librement dans ce matériau et subit de nombreux chocs au cours de son mouvement qu'on modélise globalement par une force de frottement fluide $\vec{f} = -\alpha \vec{v}$ en notant $\vec{v}$ sa vitesse et α une constante positive. Ce milieu est plongé dans un champ électrique uniforme et constant $\vec{E} = E\vec{u}_x$.

1. Déterminer l'équation différentielle vérifiée par $\vec{v}$. On négligera l'effet de la pesanteur.
2. Donner l'expression générale de la vitesse en introduisant une vitesse limite $\vec{v}_l$ et un temps caractéristique τ . On distinguera le régime transitoire et le régime permanent.

3. Sachant qu'on définit le vecteur densité de courant $\vec{j} = nq\vec{v}$ où n désigne la densité volumique de charges mobiles, donner son expression en régime permanent en fonction de n , q , α et $\vec{E}$.
4. En déduire l'expression de la conductivité électrique connaissant la loi d'Ohm locale $\vec{j} = \sigma \vec{E}$.
5. Dans la situation décrite, montrer que la force électrique est conservative et donner l'expression du potentiel en fonction de E et x .
6. En déduire que la résistance d'une portion de conducteur de section S et de longueur L soumise à une différence de potentiel U vaut $R = \frac{1}{\sigma} \frac{L}{S}$. On rappelle que l'intensité est définie comme le flux des charges positives $\vec{j}$ à travers une section S : $I = \iint \vec{j} \cdot d\vec{S}$.
7. Rappeler les hypothèses du modèle idéal de l'amplificateur opérationnel.
8. Définir le régime linéaire de l'amplificateur opérationnel.
9. En déduire la conséquence du modèle idéal de l'amplificateur opérationnel sur le régime linéaire.

montage 1 :

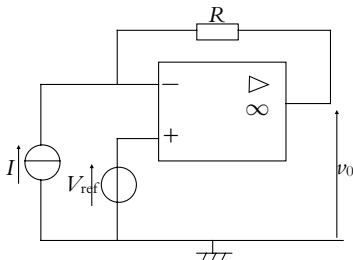


Figure 21.53

montage 2 :

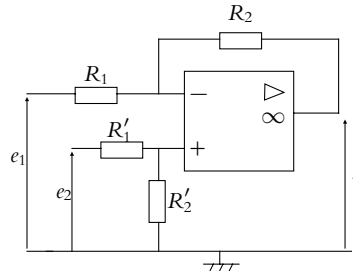


Figure 21.54

montage 3 :

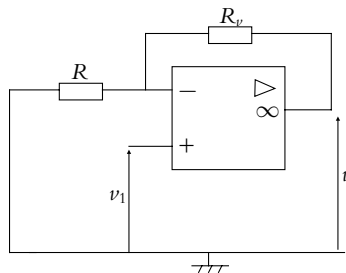


Figure 21.55

10. Dans la suite, les amplificateurs opérationnels sont supposés fonctionner en régime linéaire.

Donner l'expression de v_0 en fonction de I , V_{ref} et R pour le montage 1.

11. Exprimer s en fonction de e_1 et e_2 ainsi que des différentes valeurs de résistances pour le montage 2.

12. Que devient-elle si toutes les résistances sont identiques ? Proposer un nom à ce montage en fonction de l'opération qu'il réalise.

13. Déterminer l'expression de v en fonction de v_1 , R et R_v pour le montage 3 et qualifier ce dernier.

14. Soit le circuit ci-dessous constitué d'une capacité initialement déchargée et de deux résistances R_1 et R_2 alimenté par un signal variable $e(t)$ dont l'allure est donnée ci-dessous. Ce signal modélise une stimulation d'amplitude V_{ref} et de durée Δt .

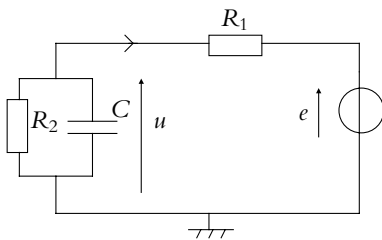


Figure 21.56

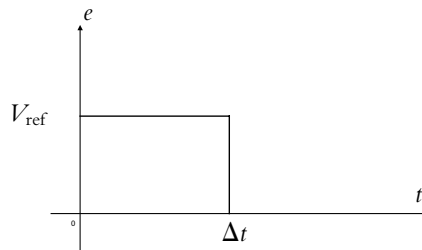


Figure 21.57

Déterminer l'équation différentielle vérifiée par la tension u aux bornes de la capacité C .

15. Donner l'expression de l'intensité traversant R_1 à l'instant 0^+ juste après le début de la stimulation.

16. Exprimer la tension $u(t)$ aux bornes de la capacité entre $t = 0$ et $t = \Delta t$. On déterminera les constantes d'intégration.

17. Donner l'expression de l'intensité traversant R_1 à l'instant Δt^- juste avant la fin de la stimulation. On suppose que la durée Δt de la stimulation est suffisante pour que le régime permanent soit atteint.

18. Exprimer l'intensité traversant R_1 en fonction du temps entre $t = 0$ et $t = \Delta t$.

19. Déterminer l'expression de $u(t)$ au-delà de $t = \Delta t$.

20. En déduire l'expression de $i(t)$ au-delà de $t = \Delta t$.

21. Les techniques de potentiel imposé à une membrane ont pour but de maintenir le potentiel membranaire d'une cellule à une valeur fixe et d'enregistrer simultanément les courants ioniques liés aux transferts d'ions à travers la membrane. On utilise donc une électrode de référence indifférente (électrode au calomel ou au chlorure d'argent) et une électrode de mesure montées en opposition. La pipette d'enregistrement est un simple tube de verre contenant une solution ionique de composition fixée par l'expérience dans lequel est placée une électrode

d'argent chlorurée. L'ensemble permet la conduction électrique entre la membrane cellulaire et le premier étage de la chaîne de mesure.

En utilisant les résultats des questions antérieures, exprimer v_0 en fonction de V_{ref} , R_f et $i(t)$.

22. Exprimer v_1 en fonction de V_{ref} et v_0 puis en fonction de R_f et $i(t)$.

23. Exprimer v en fonction de $i(t)$, R_f , R et R_v .

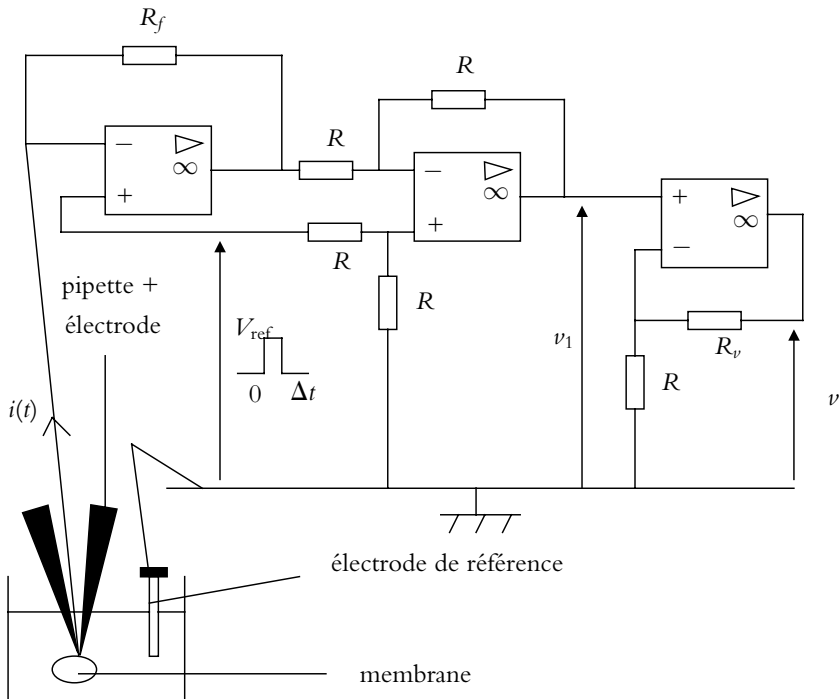


Figure 21.58

24. Électriquement la pipette est modélisable par une résistance R_p de $10 \text{ M}\Omega$. La zone de contact entre la pipette et la membrane peut être représentée par un cylindre de diamètre $d = 1,0 \text{ }\mu\text{m}$ et de hauteur $h = 2,0 \text{ }\mu\text{m}$, de conductivité $\sigma = 0,010 \text{ }\Omega^{-1}.\text{cm}^{-1}$. Exprimer la résistance R_a d'accès à la membrane en fonction de h , d et σ . Donner sa valeur numérique.

25. Il se forme par ailleurs une résistance de jonction R_s conditionnant la stabilité de la liaison pipette membrane. Cette résistance est formée par une colonne cylindrique entourant le cylindre de la zone de contact. Sa conductivité est la même, son épaisseur e est négligeable devant d et sa hauteur est h . Exprimer R_s en fonction de h , d , e et σ . Donner sa valeur numérique.

26. Quel est le montage électrique équivalent à l'association des trois résistances R_p , R_a et R_s ? Simplifier ce montage compte tenu des valeurs numériques des résistances.

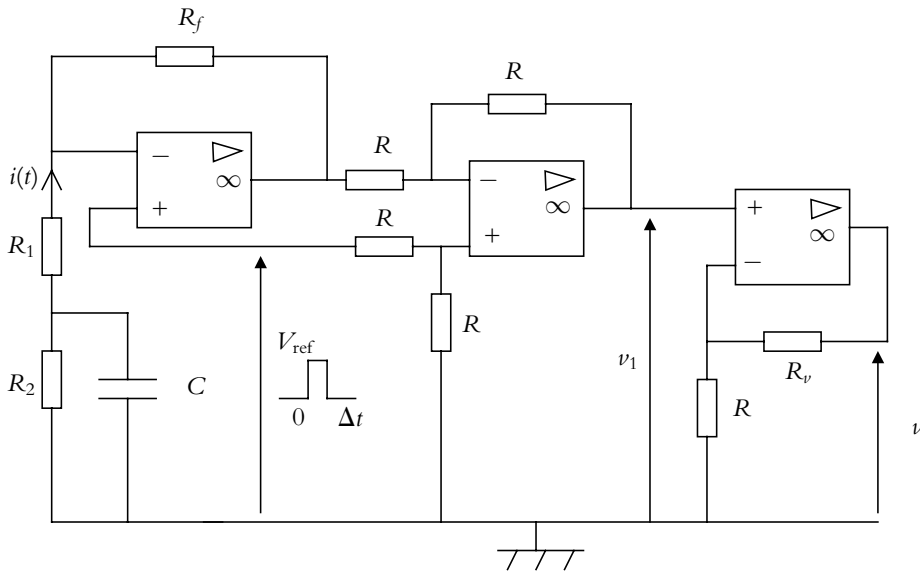


Figure 21.59

27. R_1 étant la résistance équivalente du dispositif qui vient d'être décrit et R_2 et C respectivement la résistance et la capacité de la membrane, déduire de l'enregistrement suivant les valeurs de R_1 , R_2 et C . On donne $V_{\text{ref}} = 5,0 \text{ mV}$, $R_f = 100 \text{ M}\Omega$ et $R_v = 0,0 \Omega$. L'abscisse est graduée en secondes et l'ordonnée en volts.

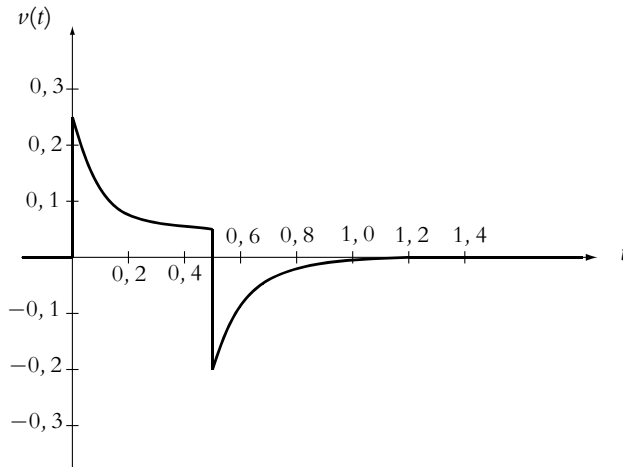


Figure 21.60

Étude expérimentale de quelques circuits à diode (PCSI)

22

1. La diode : un dipôle non linéaire

1.1 Caractéristiques externes courant tension

La caractéristique externe d'un dipôle est la courbe représentative de l'intensité en fonction de la tension pour une convention choisie, par exemple la convention récepteur. La caractéristique d'un dipôle linéaire est une droite (par exemple, celle d'un résistor)

On dit que le dipôle est non linéaire lorsque la caractéristique n'est pas une droite.

Un exemple important est étudié dans ce chapitre : *la diode à jonction*.

Son symbole est donné sur la figure 22.1.

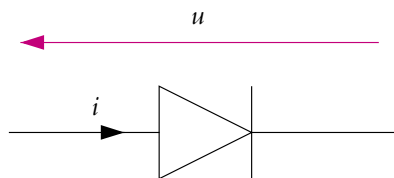


Figure 22.1 Symbole d'une diode.

1.2 Obtention expérimentale de la caractéristique de la diode

Pour tracer expérimentalement la caractéristique d'une diode à l'oscilloscope, on utilise un montage avec transformateur d'isolement déjà présenté au paragraphe 4.5 du chapitre sur l'instrumentation électrique (Cf. figure 22.2) et on obtient l'allure reproduite sur la figure 22.3.

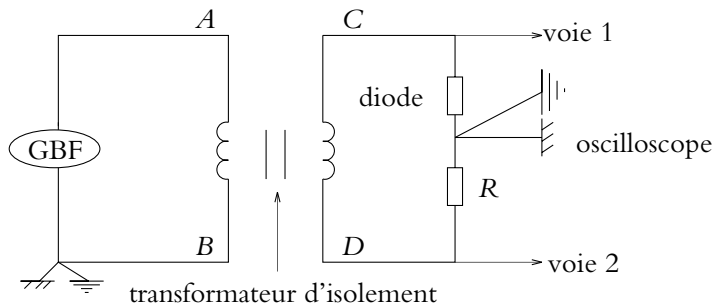


Figure 22.2 Tracé de la caractéristique d'une diode.

On remarque que la caractéristique n'est pas symétrique : on dit que la diode est un *dipôle polarisé*.

Dans la zone $u > 0$, la diode est dite *polarisée en direct* alors que dans la zone $u < 0$, elle est dite *polarisée en inverse*.

Dans la zone de polarisation directe, l'intensité croît rapidement avec la tension dès que cette dernière dépasse environ 0,5 V : on dit que la diode est *passante*. Le sens du triangle du symbole donne le sens dans lequel peut passer le courant.

En revanche, dans la zone de polarisation inverse, l'intensité est très faible (quelques nA voire quelques pA) : on dit que la diode est *bloquée ou bloquante*. Une diode agit quasiment comme un interrupteur qui laisse passer le courant uniquement si la tension est suffisamment positive.

En fait, une caractéristique plus complète de la diode est donnée sur la figure 22.4.

Lors de l'utilisation en polarisation directe il faut vérifier que l'intensité reste inférieure à une valeur maximale $i_{\max}$ qui peut varier de 20 mA pour une diode de signal utilisable en hautes fréquences à plusieurs ampères pour une diode de redressement utilisable à une fréquence voisine de celle du secteur (50 Hz).

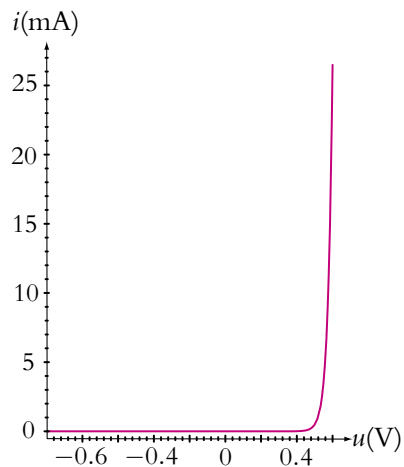


Figure 22.3 Allures de la caractéristique.

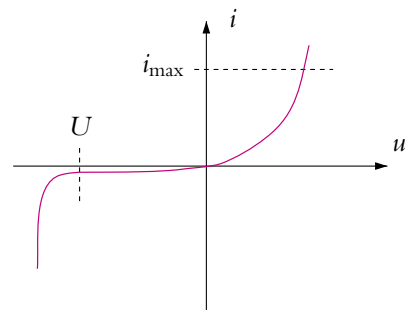


Figure 22.4 Caractéristique complète d'une diode.

Lors de l'utilisation en polarisation inverse, la tension doit rester supérieure à une valeur minimale notée $-U$.

La zone $u < -U$, appelée *zone d'avalanche*, est une zone dans laquelle la diode est détériorée. Dans cette zone, le champ électrique dû à la forte tension à laquelle la diode est soumise crée une ionisation des constituants de la diode. Les électrons arrachés acquièrent une énergie cinétique importante qui, lors de chocs avec d'autres atomes, provoquent l'arrachement de plusieurs électrons eux-mêmes accélérés par la suite. Le phénomène est à croissance exponentielle (effet « boule de neige ») et provoque la destruction du composant.

Entre les deux limites de fonctionnement précédentes, la caractéristique de la diode est exponentielle :

$$i = i_c \left(\exp \left(\frac{eu}{\eta k_B T} \right) - 1 \right) \quad (22.1)$$

où $e = 1,6 \cdot 10^{-19}$ C est la valeur absolue de la charge de l'électron, $k_B = 1,38 \cdot 10^{-23}$ J.K⁻¹ la constante de Boltzmann que l'on rencontrera dans le cours de thermodynamique, T la température en kelvin et η un coefficient variant de 1 à 2 suivant le type de diode.

1.3 Point de fonctionnement

La notion de *point de fonctionnement* est très utile pour étudier le fonctionnement des dipôles non linéaires.

On considère deux dipôles en série (Cf. figure 22.5), l'un de caractéristique $i = C_1(u)$ et l'autre $i = C_2(u)$, leur allure étant donnée par la figure 22.6. On remarque qu'un des dipôles est forcément en convention récepteur alors que l'autre est en convention générateur. L'intensité et la tension étant les mêmes pour les deux dipôles, leurs valeurs communes ne peuvent que correspondre à l'intersection des deux caractéristiques. Le point d'intersection P est appelé *point de fonctionnement*

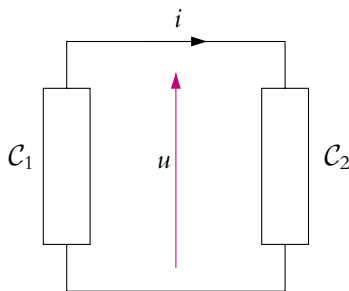


Figure 22.5 Dipôles et point de fonctionnement.

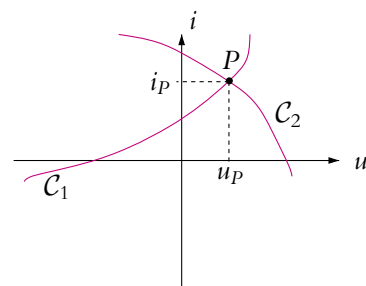


Figure 22.6 Détermination graphique d'un point de fonctionnement.

► **Remarque :** on a déjà rencontré la notion de point de fonctionnement lors de la présentation de l'alimentation stabilisée au paragraphe 2 du chapitre sur l'instrumentation électrique.

Cette notion revêt un intérêt particulier dans le cas où l'un des dipôles est un circuit linéaire représenté par son modèle de Thévenin équivalent et où l'autre dipôle est non linéaire (Cf. figure 22.7). Dans ce cas, la caractéristique du dipôle linéaire est une droite d'équation $u = e - Ri$ appelée *droite de charge*. Ses intersections avec les axes sont $i = \frac{e}{R}$ pour l'axe des ordonnées et $u = e$ pour l'axe des abscisses. Sa pente est $\frac{1}{R}$ (Cf. figure 22.8).

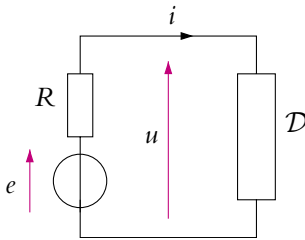


Figure 22.7 Point de fonctionnement d'un dipôle en série avec un générateur.

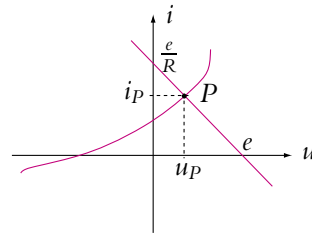


Figure 22.8 Détermination graphique du point de fonctionnement d'un dipôle en série avec un générateur.

1.4 Résistance différentielle ou dynamique

Il s'agit maintenant d'étudier le comportement du dipôle au voisinage du point de fonctionnement pour le circuit de la figure 22.7. On suppose que la f.e.m. du générateur varie d'une petite valeur de et passe donc de e à $e + de$ avec $|de| \ll |e|$. Cette variation va induire une variation di de i et du de u donc le point de fonctionnement P va se déplacer en un point voisin P' . Entre P et P' très voisins, on peut assimiler la caractéristique du dipôle à sa tangente qui a pour pente $g_d = \frac{di}{du}$, g_d étant homogène à une conductance.

On appelle *résistance dynamique* ou *différentielle* r_d , l'inverse de g_d soit :

$$r_d = \frac{du}{di} \quad (22.2)$$

On utilise cette grandeur pour modéliser le comportement d'un dipôle au voisinage d'un point de fonctionnement, ce qui est exposé dans le paragraphe suivant.

1.5 Modélisation des caractéristiques de la diode

Les caractéristiques réelles des diodes ne sont pas facilement utilisables pour étudier le fonctionnement des circuits. On va essayer de modéliser ces caractéristiques de

manière à trouver des équivalents linéaires aux dipôles suivant la zone de fonctionnement.

Si on observe de nouveau la caractéristique de la figure 22.3, en ne tenant pas compte du coude, on peut la modéliser par deux droites : l'une horizontale et l'autre proche d'une verticale. Ce modèle, qui est le plus complet des modèles représentant la diode dans sa zone de fonctionnement, est représenté sur la figure 22.9, la convention d'orientation est celle de la figure 22.1. On représente donc la caractéristique par deux droites, l'une de pente nulle et l'autre de pente $g = \frac{1}{r}$ où r est la résistance différentielle qui est très petite puisque la droite est presque verticale.

- Pour $u < u_s$, $i = 0$ A. La diode est bloquée ou bloquante et est équivalente à un interrupteur ouvert (Cf. figure 22.10).
- Pour $u \geq u_s$, $i = g(u - u_s)$ ou $u = u_s + ri$. La diode est passante et on la modélise par un générateur idéal de tension et un résistor en série (Cf. figure 22.11). D'après la convention de la figure, u_s est dans le sens opposé à i , on parlera alors de force contre électromotrice notée f.c.e.m.

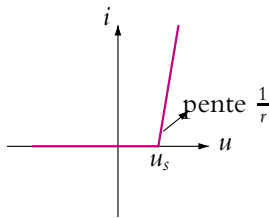


Figure 22.9 Caractéristique simplifiée d'une diode.

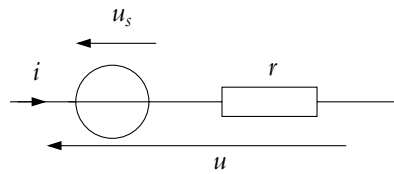


Figure 22.11 Modélisation d'une diode passante.

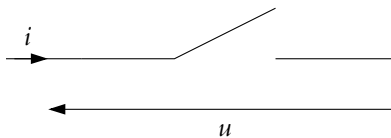


Figure 22.10 Modélisation d'une diode bloquée.

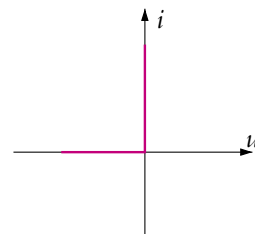


Figure 22.12 Caractéristique d'une diode idéale.

u_s est appelée *tension de seuil*. C'est à partir de cette valeur de tension que le courant devient notable. Les valeurs courantes de u_s sont 0,6 V pour les diodes au silicium et 0,4 V pour les diodes au germanium. La valeur de la résistance dynamique est de quelques ohms.

Les faibles valeurs de r et u_s permettent de donner un second modèle appelé modèle de la *diode idéale* où l'on considère $u_s = 0$ V et $r = 0$ Ω . C'est le modèle le plus simple,

qu'on utilise en général pour une première approche du fonctionnement d'un circuit. D'autre part, dans de nombreux cas d'utilisation des diodes, les tensions dans le circuit sont bien supérieures à u_s et les résistances bien supérieures à r si bien que prendre ces deux valeurs égales à 0 est une excellente approximation. Dans ce cas, on obtient la caractéristique de la figure 22.12.

- Pour $u < 0$ V, la diode est bloquée et équivalente à un interrupteur fermé.
- Pour $u = 0$ V, la diode est passante et équivalente à un court-circuit.

Les modèles intermédiaires peuvent aussi être utilisés :

- $u_s \neq 0$ V et $r = 0$ Ω ,
- $u_s = 0$ V et $r \neq 0$ Ω .

1.6 Méthode graphique pour trouver le point de fonctionnement de la diode

La méthode graphique est très efficace pour trouver la zone de fonctionnement dans le cas d'un circuit composé d'une maille (Cf. figure 22.13). On trace la caractéristique de la diode et la droite de charge $u = e - Ri$, on obtient le schéma de la figure 22.14 sur laquelle deux possibilités de point de fonctionnement sont représentées.

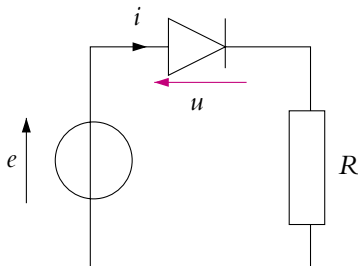


Figure 22.13 Point de fonctionnement d'une diode.

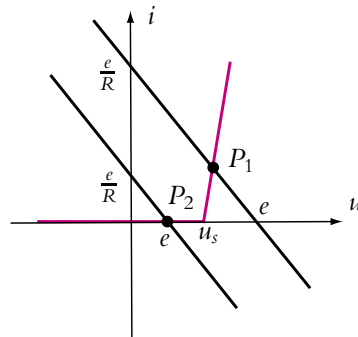


Figure 22.14 Détermination graphique du point de fonctionnement d'une diode.

1. Dans le cas où l'intersection de la droite de charge avec l'axe des abscisses a lieu pour une tension supérieure à u_s , cela signifie donc que $e \geq u_s$, la diode est passante (point P_1 de fonctionnement). On peut remplacer alors la diode par son modèle passant, ce qui donne le circuit de la figure 22.15.
2. Dans le cas où l'intersection de la droite de charge avec l'axe des abscisses a lieu pour une tension inférieure à u_s , cela signifie donc que $e \leq u_s$ et la diode est bloquée (point P_2 de fonctionnement). On obtient alors le circuit de la figure 22.16.

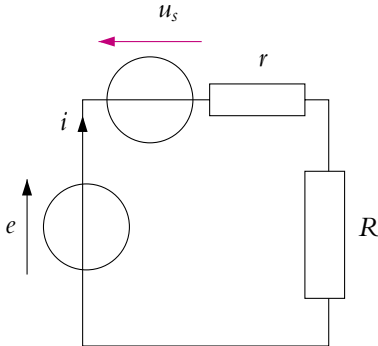


Figure 22.15 Circuit équivalent quand la diode est passante.

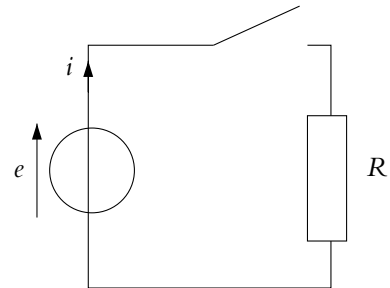


Figure 22.16 Circuit équivalent quand la diode est bloquée.

2. Redressement et filtrage

On s'intéresse à des applications de la diode à jonction qui est utilisée principalement pour transformer un signal sinusoïdal (par exemple celui délivré par le secteur) en signal continu. On peut se poser deux questions.

- Pour quelle raison a-t-on besoin de tension continue ? Tout simplement pour alimenter les appareils électroniques, comme les amplificateurs opérationnels en travaux pratiques.
- Pourquoi ne pas produire directement du courant continu ? Le courant alternatif est beaucoup plus simple à produire grâce à des alternateurs. D'autre part, comme cela a été vu dans le chapitre sur la puissance, il est préférable d'effectuer le transport sous haute tension et donc de pouvoir aisément transformer la valeur de la tension, ce qui est beaucoup plus facile pour les signaux alternatifs.

2.1 Redressement monoalternance

La première étape consiste à redresser le signal sinusoïdal, c'est-à-dire à supprimer les parties négatives (redressement monoalternance) ou à les rendre positives (redressement double-alternance). Ceci permet d'augmenter la valeur moyenne du signal qui est nulle pour un signal sinusoïdal (on rappelle que la valeur moyenne d'un signal n'est autre que sa composante continue).

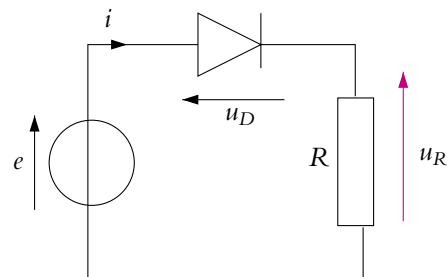


Figure 22.17 Redressement monoalternance.

On utilise le circuit de la figure 22.17 et on s'intéresse à la tension u_R aux bornes du résistor et à la tension u_D aux bornes de la diode. La f.e.m. est sinusoïdale $e = E\sqrt{2} \cos \omega t$.

a) Cas d'une alimentation de forte amplitude

On considère le cas usuel où la tension d'alimentation est celle du secteur, de valeur efficace $E = 240$ V. On observe alors les signaux suivants :

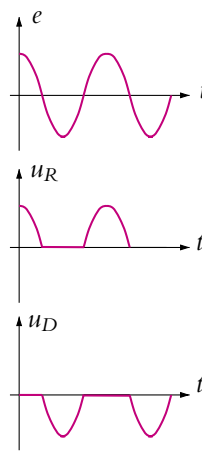


Figure 22.18 Allure des tensions pour un redressement monoalternance.

u_R se superpose parfaitement aux alternances positives de e et u_D se superpose parfaitement aux alternances négatives de e , c'est d'ailleurs pour cette raison que les courbes ont été décalées.

Étant donné la forte valeur des tensions par rapport à la tension de seuil de la diode, on va supposer la diode idéale.

Puisque $u_R = Ri$ d'après les courbes :

- $i \neq 0$ lorsque $e \geq 0$ donc la diode est passante et $u_D = 0$.
- $i = 0$ lorsque $e \leq 0$ donc la diode est bloquée.

On peut interpréter ces résultats en utilisant la méthode graphique. La droite de charge $u_D = e - Ri$ a pour intersection avec les axes $u_D = e$ et $i = \frac{e}{R}$ et sa pente est $\frac{1}{R}$ (Cf. figure 22.19). Au cours du temps, e varie tandis que R reste constante : la droite se déplace donc parallèlement à elle-même entre deux droites extrêmes (en pointillés) d'intersection avec l'axe des abscisses : $e_{\max} = E\sqrt{2}$ (droite D_1) et $e_{\min} = -E\sqrt{2}$ (droite D_2).

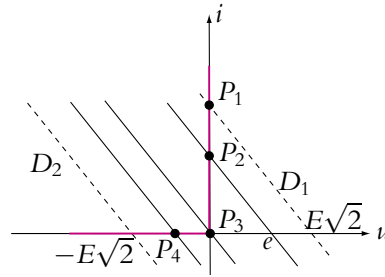


Figure 22.19 Évolution du point de fonctionnement lors d'un redressement monoalternance.

1. À l'instant $t=0$, la droite de charge est confondue avec D_1 , le point de fonctionnement est P_1 . Ainsi on lit sur la caractéristique que $u_D=0$ V d'où $u_R = e - u_D = e$.
2. Lorsque le temps augmente, e diminue et la droite de charge se déplace vers la gauche (point de polarisation P_2), la tension u_D reste nulle tant que e reste positif jusqu'au point de polarisation P_3 . Ainsi pour $e > 0$ V, $u_D = 0$ V et $u_R = e$.
3. e devient négative et la droite de charge continue à se déplacer vers la gauche (point de polarisation P_4) jusqu'à atteindre la droite D_2 . Dans toute cette zone, l'intensité est nulle d'après la position du point de polarisation. On en déduit que $u_R = Ri = 0$ V et que $u_D = e$.
4. e recommence alors à croître, la droite se déplace vers la droite. Tant que e reste négatif (point P_4), $u_R = Ri = 0$ V et $u_D = e$, puis dès que e devient positif (point P_2), $u_D = 0$ V et $u_R = e$.

La tension aux bornes de R ne comporte que des alternances positives, on dit qu'elle a été redressée. Comme il n'y a qu'une alternance sur 2, on parle de *redressement monoalternance*.

On peut calculer la valeur moyenne du signal ainsi redressé :

$$\langle u_R \rangle = \frac{1}{T} \int_0^T u_R(t) dt$$

Or $u_R(t) = E\sqrt{2} \cos \omega t$ pour $0 < t < T/4$ et $3T/4 < t < T$ et est nulle sinon, d'où :

$$\langle u_R \rangle = \frac{1}{T} \int_0^{T/4} E\sqrt{2} \cos(\omega t) dt + \frac{1}{T} \int_{3T/4}^T E\sqrt{2} \cos(\omega t) dt$$

soit

$$\langle u_R \rangle = \frac{E\sqrt{2}}{\omega T} \left([\sin \omega t]_0^{T/4} + [\sin \omega t]_{3T/4}^T \right) = \frac{E\sqrt{2}}{\pi}$$

On obtient donc bien une composante continue positive.

b) Cas d'une alimentation de faible amplitude

Lorsque la tension d'alimentation a une amplitude comparable à celle de la tension de seuil (par exemple 2 V), on voit apparaître des déformations du signal u_R dues au fait que la diode n'est pas idéale. On observe les signaux de la figure 22.20 :

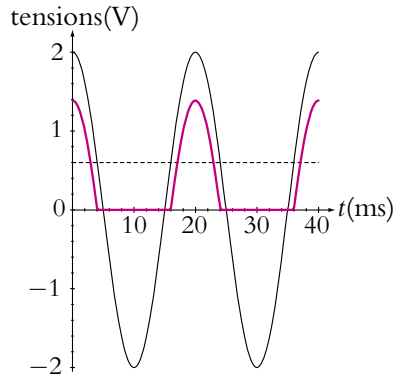


Figure 22.20 Redressement monoalternance d'une faible tension (trait plein fin : e , trait plein épais : u_R , trait pointillé : u_s).

On remarque donc que tant que e reste inférieure à u_s , u_R est nulle et que, quand u_R n'est pas nulle, elle ne se superpose pas à e comme dans le cas précédent.

On va interpréter ces phénomènes en tenant compte de la tension de seuil $u_s = 0,6 \text{ V}$ et de la résistance dynamique r . On utilise le même raisonnement graphique que précédemment mais sur la figure 22.21.

Dans la zone passante de la diode (point de polarisation P_2), on obtient le schéma équivalent du circuit de la figure 22.15 étudié dans le paragraphe de présentation de la méthode graphique. Dans ce cas :

$$u_R = Ri = R \frac{e - u_s}{R + r} < e$$

La diode est passante tant que $e > u_s$ puis i est nulle et u_R aussi. On obtient alors le signal de la figure 22.20 pour u_R . La valeur maximale de u_R n'est plus $E\sqrt{2}$ mais $R \frac{E\sqrt{2} - u_s}{R + r}$.

Dans les conditions usuelles, par exemple avec la tension du secteur $E = 240 \text{ V}$, ces écarts ne sont pas visibles car $u_s \ll E$ et l'on retrouve les signaux présentés dans le paragraphe précédent.

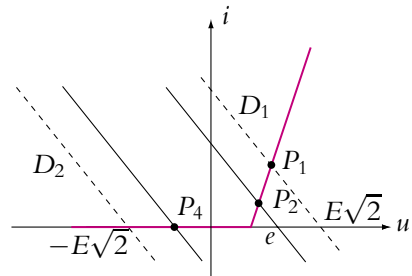


Figure 22.21 Évolution du point de fonctionnement lors du redressement monoalternance d'une faible tension.

2.2 Redressement double-alternance

Pour augmenter la valeur moyenne et se rapprocher encore plus d'un signal continu, il faut pouvoir redresser la deuxième alternance. Pour cela, on utilise un circuit à quatre diodes appelé *pont de Graetz* (Cf. figure 22.22).

On obtient le signal de la figure 22.23 qui porte le nom de *signal redressé double-alternance*. On remarque que quel que soit le signe de la f.e.m. e , u_R et donc i sont positifs. Le courant passe *a priori* par la diode D_2 sur une demi-alternance et par la diode D_4 pendant l'autre demi-alternance. Ces deux diodes ne peuvent pas être passantes simultanément car sinon le générateur serait court-circuité. Il en est de même des diodes D_1 et D_3 .

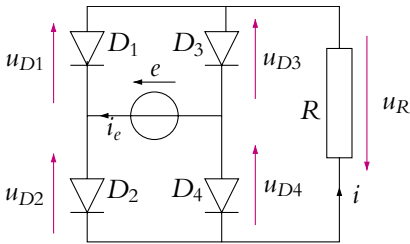


Figure 22.22 Redressement double alternance : pont de Graetz.

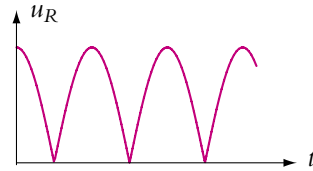


Figure 22.23 Allure d'un signal redressé double alternance.

Pour interpréter la forme de u_R , on suppose que la f.e.m. a la même forme que dans le redressement monoalternance ($e(t) = E\sqrt{2} \cos \omega t$) et que les diodes sont idéales.

Des lois des mailles permettent d'écrire les relations suivantes :

$$u_R = -u_{D1} - u_{D2} \quad (22.3)$$

$$u_R = -u_{D3} - u_{D4} \quad (22.4)$$

$$e = u_{D2} - u_{D4} \quad (22.5)$$

$$e = u_{D3} - u_{D1} \quad (22.6)$$

À $t = 0$, e est positive donc le courant i_e est positif, il ne peut passer que dans la diode D_2 puisque D_1 est polarisée en inverse. Dans ce cas, la diode D_2 est passante et $u_{D2} = 0$ V.

On déduit de l'équation (22.5) que $u_{D4} = -e$ est négative puisque e est positive et donc que la diode D_4 est bloquée. Le courant i_e circule dans R et $i = i_e$, la tension u_R est donc positive.

Pour revenir, le courant i doit passer par l'une des diodes D_1 ou D_3 . Si D_1 était passante, $u_{D1} = 0$ V et d'après la relation (22.3), u_R serait nulle donc i et i_e également, ce qui n'est pas le cas. Ainsi D_1 ne peut être que bloquée et D_3 passante ($u_{D3} = 0$ V). Alors, d'après la relation (22.4), $u_R = -u_{D4}$ et avec $u_{D4} = -e$:

$$u_R = e = E\sqrt{2} \cos \omega t \quad \text{positif tant que } e > 0 \text{ V}$$

Lorsque e devient négatif, par un raisonnement analogue au précédent, on trouve que D_2 et D_3 sont bloquées alors que D_1 et D_4 sont passantes. Le courant i_e change de sens dans le générateur mais grâce aux diodes, le courant i a toujours le même sens. Le calcul donne alors :

$$u_R = -e = -E\sqrt{2} \cos \omega t \quad \text{positif tant que } e < 0 \text{ V}$$

Puisqu'il y a deux fois plus d'arches que dans le redressement monoalternance, la valeur moyenne est deux fois plus élevée que pour le redressement double-alternance, ce que le lecteur pourra vérifier en calculant l'intégrale :

$$\langle u_R \rangle = \frac{1}{T} \int_0^T |E\sqrt{2} \cos \omega t| dt = \frac{2\sqrt{2}E}{\pi}$$

dont on a déjà parlé au paragraphe sur l'utilisation des multimètres.

2.3 Filtrage avec un condensateur

Pour obtenir un signal continu, il faut supprimer l'ondulation du signal redressé. Pour cela, on utilise un condensateur qui stocke de l'énergie quand la diode est passante et la restitue au résistor quand la diode est bloquée. Le circuit utilisé est représenté sur la figure 22.24.

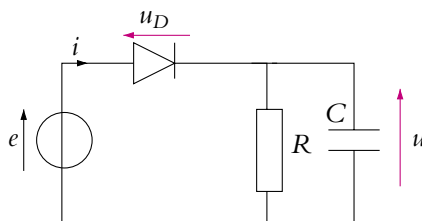


Figure 22.24 Filtrage d'un signal redressé monoalternance.

On appelle $\tau = RC$, la constante de temps du circuit R, C .

Le générateur délivre une f.e.m. $e(t) = E\sqrt{2} \cos \omega t$ de période T . On a représenté les signaux obtenus sur les figures 22.25, 22.26, 22.27, 22.28 pour quatre valeurs de τ différentes et croissantes. La f.e.m. est représentée en trait noir, alors que la tension u l'est en trait rouge.

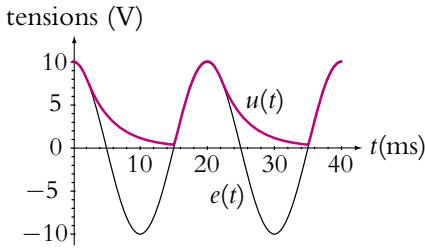


Figure 22.25 Allure d'un signal redressé monoalternance et filtré pour $\tau = T/5$.

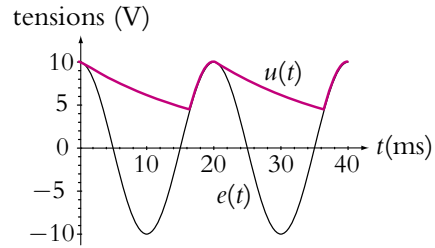


Figure 22.26 Allure d'un signal redressé monoalternance et filtré pour $\tau = T$.

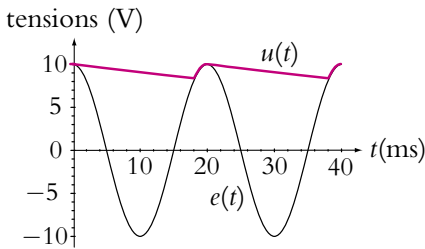


Figure 22.27 Allure d'un signal redressé monoalternance et filtré pour $\tau = 5T$.

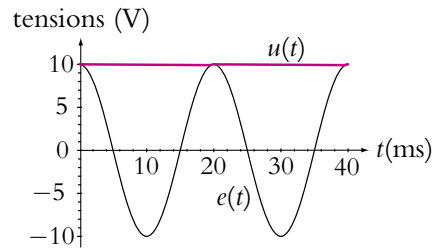


Figure 22.28 Allure d'un signal redressé monoalternance et filtré pour $\tau = 100T$.

On constate que pour une valeur de τ grande devant la période T (Cf. figure 22.28), u est quasiment une tension continue. Si on analyse une période des autres figures, il apparaît nettement deux phases :

- une première phase pour laquelle $u = e$,
- une deuxième phase qui ressemble à un régime transitoire de circuit R, C .

Dans la première phase, puisque $u = e$, la diode est passante et se comporte en court-circuit. Dans la deuxième phase, puisque $u \neq e$, la diode est bloquée et le circuit R, C est isolé du générateur.

On va étudier ces différentes phases en supposant la diode idéale, en raison de l'ordre de grandeur de $E \gg u_s$.

À l'instant initial, e est positive et impose un courant positif : la diode est passante. Dans ce cas, u_D est nulle et $u = e$. Pour déterminer l'instant auquel la diode se bloque, il faut déterminer l'instant d'annulation de l'intensité i . Tant que la diode est passante, elle est équivalente à un simple fil, le montage est donc linéaire et fonctionne en régime sinusoïdal forcé. Les tensions e et u sont sinusoïdales, on peut utiliser les complexes et :

$$i = \underline{Y} \underline{e} = \left(\frac{1}{R} + jC\omega \right) \underline{e}$$

si l'on note $\underline{i} = I\sqrt{2} \exp(j\omega t + j\varphi)$ et $\underline{e} = E\sqrt{2} \exp(j\omega t)$. Avec l'équation précédente, on détermine le déphasage $\varphi = \arctan RC\omega > 0$ (car $\cos \varphi > 0$). Ainsi l'intensité réelle $i(t) = I\sqrt{2} \cos(\omega t + \varphi)$ est en avance de φ sur e et devient nulle à l'instant $t_1 = \frac{T}{4} - \frac{\varphi}{\omega}$. Sur l'intervalle $[0, t_1]$, $u = e = E\sqrt{2} \cos(\omega t)$.

Ensuite l'intensité s'annule et deviendrait négative si la diode ne se bloquait pas. La diode bloquée isole le circuit R, C du générateur (donc la tension u n'est plus sinusoïdale) et elle vérifie l'équation différentielle :

$$\frac{du}{dt} + \frac{u}{\tau} = 0 \quad \text{avec} \quad \tau = RC$$

dont la solution est $u = A \exp\left(-\frac{t}{\tau}\right)$. On détermine la constante A à l'instant t_1 par continuité de la tension u aux bornes du condensateur, soit :

$$u = E\sqrt{2} \cos(\omega t_1) \exp\left(-\frac{(t - t_1)}{\tau}\right)$$

La tension u décroît exponentiellement en restant positive. Pendant ce temps, la f.e.m. e , après avoir été négative, redevient positive si bien qu'à l'instant t_2 , on a de nouveau $u = e$, jusqu'à ce que i s'annule de nouveau à $t_3 = t_1 + T$.

On récapitule les résultats :

$$\left\{ \begin{array}{lll} 0 \leq t \leq t_1 & u(t) = e(t) = E\sqrt{2} \cos \omega t & i > 0 \text{ A} \\ t_1 \leq t \leq t_2 & u(t) = E\sqrt{2} \cos(\omega t_1) \exp\left(-\frac{(t - t_1)}{\tau}\right) & i = 0 \text{ A} \\ t_2 \leq t \leq t_1 + T & u(t) = e(t) = E\sqrt{2} \cos \omega t & i > 0 \text{ A} \end{array} \right. \quad (22.7)$$

La diode se débloque et u reste égale à e tant que i reste positif.

► Remarques

1. Si on tient compte de la tension de seuil u_s de la diode, les courbes précédentes ont la même allure en remplaçant $e(t)$ par $e(t) - u_s$ quand la diode est passante. En particulier, si $E\sqrt{2} < u_s$, la diode est toujours bloquée.
2. le résultat ne peut être que meilleur si on filtre une tension double-alternance au lieu d'une tension monoalternance.

2.4 Détecteur de crête

Un détecteur de crête est un circuit qui permet de récupérer l'enveloppe d'un signal électrique. Pour cela, on utilise le même filtre que pour le filtrage d'un courant redressé c'est-à-dire l'association d'une diode et d'un circuit R, C (Cf. figure 22.24).

En revanche, il est nécessaire d'utiliser ici une diode rapide (c'est-à-dire qui fonctionne à haute fréquence). Le signal dont on veut détecter l'enveloppe est noté V_e , et le signal obtenu après le détecteur de crête est V_s (Cf. figure 22.29).

Pour reconstituer l'enveloppe de V_e , il est nécessaire de bien choisir la constante de temps $\tau = RC$. On nomme $f_m = 1/T_m$, la fréquence et la période de V_s et $F_M = 1/T_M$ la fréquence et la période des ondulations de V_e (Cf. figure 22.29). Sur la figure, on a pris T_m et T_M très proches pour une meilleure compréhension mais dans la réalité elles sont très différentes. Il faut que τ soit grand devant T_M pour ne pas suivre l'ondulation (Cf. figure 22.30) et que τ soit petit devant T_m pour bien suivre toutes les crêtes. On a représenté sur la figure 22.30 le cas où $\tau \ll T_M$ (en gras) et le cas où $\tau \gg T_M$ (en pointillés).

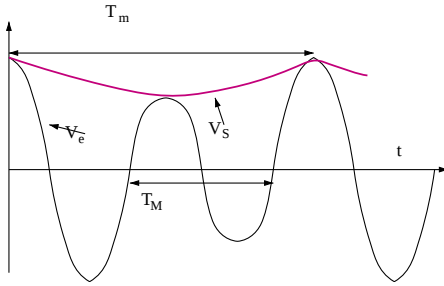


Figure 22.29 Principe d'un détecteur de crête.

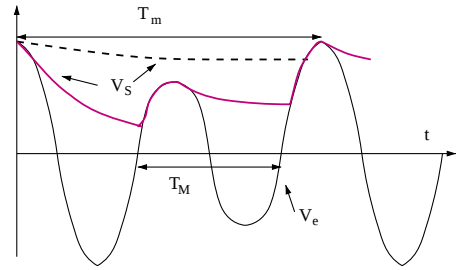


Figure 22.30 Influence du temps caractéristique du filtre R, C sur la détection de crête.

Pour obtenir une bonne détection de crête, le filtre R, C doit vérifier :

$$f_m \ll \frac{1}{RC} \ll f_M$$

3. Enrichissement du spectre par un dipôle non linéaire

Contrairement à ce qui a été vu avec les dipôles linéaires qui ne peuvent qu'appauvrir le spectre d'un signal, les dipôles non linéaires peuvent faire apparaître des fréquences qui n'existent pas dans le spectre du signal d'entrée : on dit qu'il y a *enrichissement du spectre*.

Si on prend l'exemple de la diode dans le cas du redressement monoalternance, le signal d'entrée $e(t) = E\sqrt{2}\cos\omega t$ est sinusoïdal de fréquence $f = \frac{\omega}{2\pi}$. Le signal redressé v_s contient une composante continue donc une composante de fréquence nulle et des discontinuités sont apparues. Or des discontinuités dans le signal correspondent à des harmoniques de fréquences plus élevées que le signal. Dans le cas du signal redressé monoalternance, on obtient un spectre qui a l'allure de celui de la

figure 22.31. Il contient la composante du signal e mais est plus riche que le spectre de e qui ne contient qu'une composante de fréquence f .

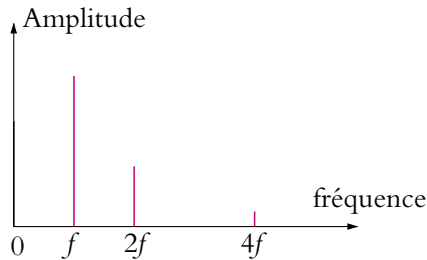


Figure 22.31 Spectre d'un signal redressé monoalternance.

4. Multiplieurs analogiques

Les multiplieurs sont des quadripôles qui permettent de multiplier deux signaux (Cf. figure 22.32). Le signal de sortie est :

$$v_s = k v_1 v_2 \quad (22.8)$$

où k est une constante caractéristique du multiplieur, s'exprimant en V^{-1} .

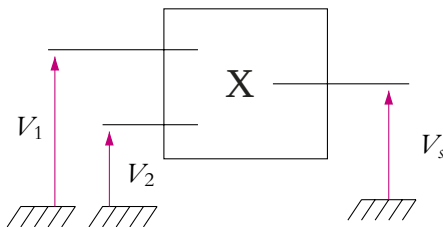


Figure 22.32 Symbole d'un multiplieur.

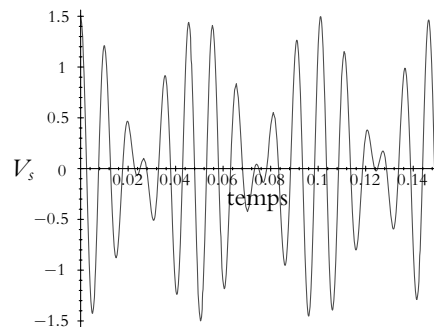


Figure 22.33 Signal obtenu par multiplication de deux sinusoïdes.

4.1 Multiplication de deux signaux sinusoïdaux

On envoie à l'entrée du multiplieur un signal de fréquence f_m :

$$v_1 = s_m(t) = A_m \cos(2\pi f_m t)$$

et un signal de fréquence f_p :

$$v_2 = s_p(t) = A_p \cos(2\pi f_p t)$$

On considère $f_p > f_m$. Le signal obtenu en sortie est :

$$v_s(t) = kA_m A_p \cos(2\pi f_m t) \cos(2\pi f_p t)$$

dont l'allure est représentée sur la figure 22.33.

Pour obtenir le spectre de ce signal, on linéarise l'expression, ce qui donne :

$$v_s(t) = \frac{k}{2} A_m A_p \left(\cos(2\pi(f_m + f_p)t) + \cos(2\pi(f_p - f_m)t) \right)$$

Le signal v_s est donc la somme de deux signaux sinusoïdaux de fréquence $f_p - f_m$ et $f_p + f_m$ et d'amplitude $\frac{k}{2} A_m A_p$.

4.2 Modulation et démodulation d'amplitude

a) Modulation d'amplitude

Un signal sinusoïdal est dit modulé en amplitude si son amplitude dépend du temps et est de la forme : $s(t) = A(t) \cos(2\pi f t)$ Pour une modulation d'amplitude, on utilise aussi une multiplication de signaux. Le signal v_1 est $v_1 = s_m(t) = A_m(1 + m \cos(2\pi f_m t))$ tandis que v_2 est inchangé. On obtient alors v_s :

$$v_s(t) = kA_m A_p (1 + m \cos(2\pi f_m t)) \cos(2\pi f_p t)$$

que l'on peut écrire sous la forme suivante en linéarisant l'expression :

$$v_s(t) = kA_m A_p \cos(2\pi f_p t) + \frac{km}{2} A_m A_p \left(\cos(2\pi(f_m + f_p)t) + \cos(2\pi(f_p - f_m)t) \right)$$

Dans la pratique, le signal s_p appelé *porteuse* a une fréquence beaucoup plus élevée que le signal s_m appelé *signal modulant*. Le paramètre m est appelé *taux de modulation* et influence la forme du signal. Le signal modulé v_s est représenté sur les figures 22.34 et 22.35 dans les deux cas $m < 1$ et $m > 1$. On parle donc de modulation d'amplitude car le signal v_s a une amplitude déterminée par le signal basse fréquence s_m .

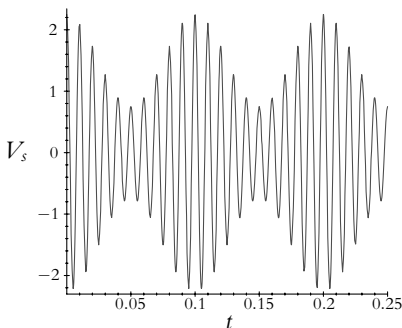


Figure 22.34 Modulation avec $m < 1$

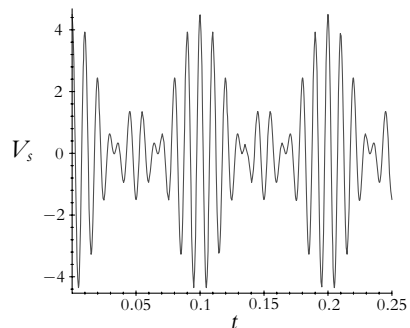


Figure 22.35 Modulation avec $m > 1$

D'après l'expression de v_s , on se rend compte que le spectre d'un signal modulé en amplitude est composé de trois fréquences :

1. une composante de fréquence $f_p - f_m$ et d'amplitude $\frac{kA_m A_p}{2}$,
2. une composante de fréquence f_p et d'amplitude $kA_m A_p$,
3. une composante de fréquence $f_p + f_m$ et d'amplitude $\frac{kA_m A_p}{2}$.

Quel est l'intérêt de moduler un signal en amplitude pour le transmettre par radio ? En fait il y a deux raisons. Tout d'abord, le domaine de fréquence des signaux audibles se situe environ entre 20 Hz et 20 kHz. Toute personne désirant émettre un signal radio choisira donc la même bande de fréquences et toutes les émissions seront brouillées. Pour utiliser la modulation d'amplitude, on multiplie le signal à envoyer par un signal sinusoïdal de plus haute fréquence $f_p \simeq 100$ kHz par exemple. Le signal à transmettre module donc l'amplitude de cette porteuse. L'avantage est que, grâce à la modulation, on a décalé la bande de fréquences transmises en hautes fréquences autour de f_p . Il suffit donc de choisir une fréquence de porteuse par radio pour éviter que les signaux transmis soient à la même fréquence.

La deuxième raison est relative à la taille des antennes d'émission et de réception. Une antenne doit avoir une taille de l'ordre de grandeur de la longueur d'onde. Or, dans le domaine des fréquences audibles, cela donnerait des antennes de plusieurs dizaines de kilomètres de haut difficilement utilisables. Le fait de transposer les signaux en hautes fréquences permet de diminuer la longueur d'onde et de ramener les antennes à des tailles raisonnables. Cependant, si l'on calcule la longueur d'onde λ correspondant aux fréquences typiques d'émission en modulation d'amplitude, soit 100 kHz (c'est ce qu'on appelle les Grandes Ondes), on trouve : $\lambda \simeq 3$ km. Contrairement au cas de la Modulation de Fréquences (fréquence de 100 MHz), on ne peut donc réaliser des antennes de la taille de la longueur d'onde, et on pallie ce problème en augmentant la puissance des émetteurs. Ainsi "France Inter", pour son canal Grandes Ondes, dispose d'un seul émetteur très puissant au centre de la France.

Un inconvénient de la modulation d'amplitude est sa forte sensibilité aux parasites. En effet, les parasites perturbent l'amplitude des signaux. Or c'est dans l'enveloppe d'un signal modulé que se trouve l'information à transmettre. C'est l'une des raisons pour lesquelles on préfère utiliser la modulation de fréquence.

b) Démodulation du signal

Pour récupérer l'information $s_m(t)$ contenue dans un signal modulé en amplitude, il faut récupérer l'enveloppe. Dans le cas où $m < 1$, on utilise alors un détecteur de crête décrit au paragraphe 2.4 avec la condition :

$$f_m \ll \frac{1}{RC} \ll f_p$$

Dans le cas où $m > 1$, le détecteur de crête ne donne pas un résultat satisfaisant. En effet, dans ce cas, les enveloppes se croisent (Cf. figure 22.36) et un détecteur de crête récupérera l'enveloppe supérieure qui est différente du signal sinusoïdal modulant.

On doit donc utiliser un autre procédé : la *démodulation synchrone*. Pour cela on multiplie $v_s(t)$ par la porteuse de fréquence f_p . Le lecteur pourra montrer que le spectre du signal obtenu est constitué de plusieurs composantes de fréquences : f_m , $f_m + 2f_p$, $2f_p - f_m$, $2f_p$ et d'une composante continue. La composante continue et la composante de fréquence f_m ont alors exactement l'amplitude qu'elles ont dans le signal $s_m(t)$. Il suffit alors d'envoyer ce signal sur un filtre passe-bas pour ne garder que le signal basse fréquence (f_m) et la composante continue.

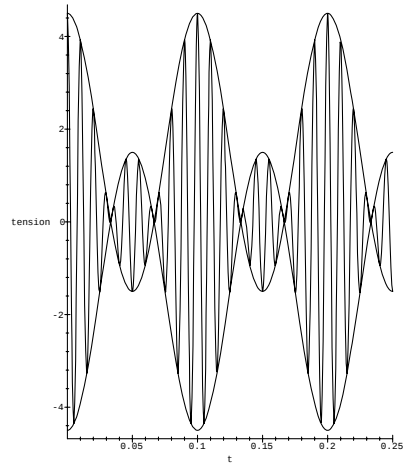


Figure 22.36 Limitations de la démodulation par détecteur de crête.

EXERCICES

A. Applications directes du cours

1. Influence d'une diode sur un circuit

La caractéristique d'une diode est donnée de façon approchée par :

$$\begin{cases} I = 0 & \text{pour } U \leq U_0 \\ I = \frac{U - U_0}{R_d}, & \text{pour } U \geq U_0 \end{cases} \quad (22.9)$$

On place cette diode en série avec une alimentation $E = 6 \text{ V}$ et une résistance $R = 200 \, \Omega$.

1. Quel est le point de fonctionnement du circuit ?
2. Calculer le facteur de régulation $f_0 = \frac{\Delta E}{\Delta U}$ et le taux d'ondulation $\frac{\Delta U}{U_0}$ si la tension de l'alimentation est un signal sinusoïdal de tension crête à crête de 2 V et de valeur moyenne 6 V.
3. Quelles sont les valeurs que doit prendre une résistance R_C placée en parallèle aux bornes de la diode pour que la diode soit passante et que son courant maximal soit de 20 mA ?

2. Théorème de Millman et dipôle non linéaire

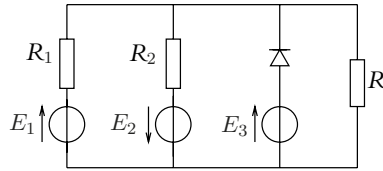


Figure 22.37

La diode admet pour tension seuil U_S et comme résistance interne R_d . Déterminer l'intensité parcourant la résistance R en appliquant le théorème de Millman.

3. Redressement sans seuil

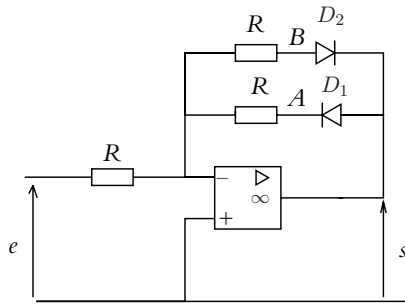


Figure 22.38

1. Rappeler le modèle simplifié de la diode en tenant compte de la tension seuil U_S et de la résistance R_d .
2. Préciser les conditions pour que chacune des diodes soit bloquée en fonction de s , U_S et des tensions aux points A et B .
3. On se place dans le cas où la diode reliée à A est bloquée et l'autre passante.
 - a) Déterminer l'expression des tensions aux points A et B ainsi que de s en fonction des résistances, de la tension seuil et de e .
 - b) Définir les conditions sur e pour se trouver dans cette situation.
4. Mêmes questions dans le cas où la diode reliée à B est bloquée et l'autre passante.
5. Est-il nécessaire d'envisager les autres cas ? Justifier.
6. En déduire le tracé des tensions aux points A et B ainsi que de s en fonction du temps.

4. Redressement

Les deux générateurs délivrent une tension sinusoïdale $e = E \cos \omega t$. Les diodes sont supposées idéales. On observe expérimentalement les tensions u_1 et u_2 représentées sur les figures 22.39.

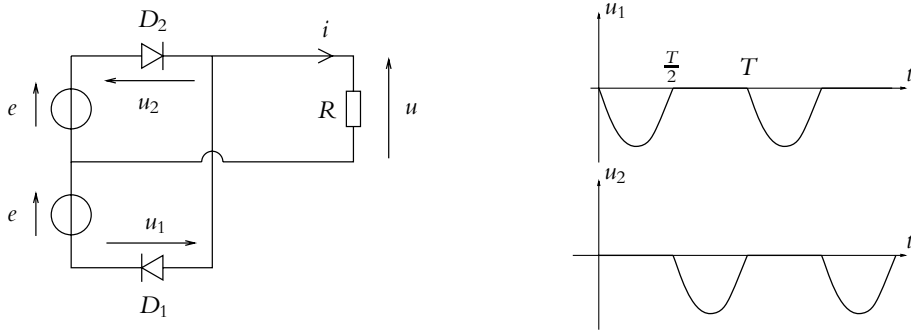


Figure 22.39

1. Déterminer deux relations entre les tensions e , u_1 , u_2 et u .
2. Déterminer les comportements des diodes D_1 et D_2 d'après l'allure des signaux u_1 et u_2 .
3. En déduire l'intensité i dans la résistance sachant que les deux générateurs délivrent une tension sinusoïdale $e = E \cos \omega t$. Les diodes sont supposées idéales. Quels sont les avantages et les inconvénients d'un tel montage par rapport à un pont de Graetz ?

B. Exercices et problèmes

1. Alimentation à découpage, d'après CCP PSI 2001

On considère une alimentation à découpage qui constitue une application des diodes et des transistors de puissance à la conversion électronique de puissance.

On utilise le transistor en interrupteur commandé. Sa caractéristique courant - tension a alors l'allure suivante :

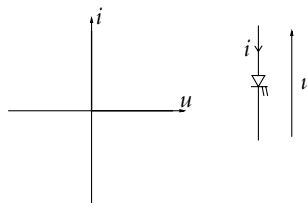


Figure 22.40

Aucune connaissance supplémentaire sur le fonctionnement du transistor n'est requise pour ce problème.

1. Caractéristique d'une diode en interrupteur commandé :

a) Tracer l'allure de la caractéristique courant - tension d'une diode qu'on suppose idéale. On précisera les orientations choisies pour le courant et la tension.

b) Proposer un montage permettant de tracer la caractéristique sur un écran d'oscilloscope sachant que tous les appareils sont reliés à la Terre et qu'on dispose d'un transformateur d'isolement, d'un oscilloscope, d'un GBF et d'une résistance R . On notera X et Y les deux voies de

l'oscilloscope. Préciser la procédure utilisée notamment le mode de visualisation de l'oscilloscope et les réglages du signal du GBF.

c) Quel est le rôle du transformateur d'isolement ?

d) On écrit la tension au secondaire du transformateur $v = V \cos \omega t$. Tracer, en les justifiant, la tension u aux bornes de la diode et la tension w aux bornes de la résistance.

e) Que se passe-t-il si l'expérimentateur veut faire une mesure en régime continu sans changer de montage ? Donner notamment le courant à travers R .

2. On ne s'intéresse qu'au fonctionnement en régime périodique (de période T) de l'alimentation à découpage suivante :

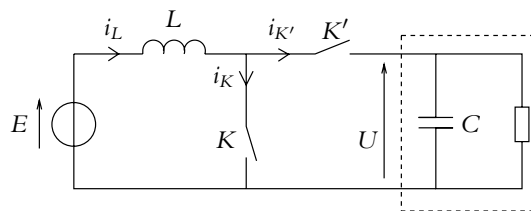


Figure 22.41

Les interrupteurs K et K' sont des interrupteurs commandés dont la séquence de commande est la suivante : K est fermé et K' ouvert pour $0 \leq t \leq \alpha T$, K est ouvert et K' fermé pour $\alpha T \leq t \leq T$. On suppose dans un premier temps que l'association en parallèle de R et C se comporte comme une source idéale de tension de force électromotrice $U = E'$. On se place dans l'hypothèse où le courant dans la bobine d'inductance L ne s'annule jamais. On prend $E = 50 \text{ V}$ et $T = 50 \mu\text{s}$.

a) Déterminer les expressions des courants i_L , i_K et $i_{K'}$ en fonction du temps sur une période. On note I_m la valeur minimale de i_L .

b) Représenter l'allure de ces courants en fonction du temps.

c) Déterminer la valeur de E' en fonction de E et α .

d) On règle le rapport cyclique α à 0,6. Pour cette valeur du rapport cyclique, on accepte une ondulation maximale $\Delta i_{L, \max}$ de 0,3 A en notant $\Delta i_L = I_M - I_m$ l'ondulation avec I_M la valeur maximale de i_L . Déterminer l'expression littérale puis la valeur numérique de la valeur minimale de L .

e) La puissance fournie par la source de tension de force électromotrice E est $P_g = 150 \text{ W}$. Établir l'expression de P_g en fonction de I_m , E , α , T et L .

f) En déduire, pour la valeur minimale de L , les valeurs numériques de I_m et I_M .

g) Tracer les portions de caractéristiques courant - tension décrites par chaque interrupteur sur une période.

h) En déduire la nature (diode ou transistor) des interrupteurs K et K' .

i) Déterminer la valeur moyenne V_0 de la tension aux bornes de K en fonction de E .

3. On garde les conditions $\alpha = 0,6$ et $P_g = 150 \text{ W}$ mais on considère que la tension aux bornes de l'association en parallèle de R et C n'est pas constante mais une fonction périodique présentant une légère ondulation. On fait l'hypothèse que cela ne modifie pas les évolutions temporelles de i_L , i_K et $i_{K'}$.

- Déterminer littéralement et numériquement l'intensité moyenne I_C du courant dans la capacité C .
- Même question pour l'intensité moyenne I_R du courant dans la résistance R . On donnera l'expression littérale en fonction de E , P_g et α .
- Calculer la puissance moyenne P_C dissipée dans la capacité.
- Même question pour la puissance moyenne P_R dissipée dans la résistance.
- Que peut-on en conclure à propos de l'hypothèse faite auparavant sur E' ?

2. Étude d'un hacheur

Le hacheur permet d'obtenir une tension continue réglable à partir d'une alimentation constante. Cette tension variable est appliquée à un moteur afin de faire varier sa vitesse.

Le hacheur est constitué :

- d'un interrupteur H qui :
 - ne laisse passer le courant i_H que dans le sens indiqué,
 - se ferme et s'ouvre périodiquement avec une période $T = 1 \text{ ms}$: il est fermé pour $0 < t < \alpha T$ avec $0 < \alpha < 1$ et il est ouvert le reste de la période.
- d'une diode idéale ;
- d'une bobine d'inductance $L = 40 \text{ mH}$.

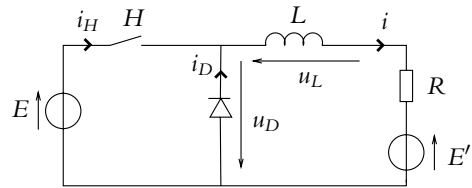


Figure 22.42

L'alimentation est une source de tension continue $E = 100 \text{ V}$. La charge est un moteur modélisé par une résistance $R = 1 \Omega$ et une source idéale de tension de f.c.e.m. $E' = 60 \text{ V}$. On se trouve dans des conditions telles que le courant $i(t)$ ne s'annule jamais.

- L'interrupteur H est fermé ($0 < t < \alpha T$).
 - La diode est-elle passante ou bloquée ? En déduire le schéma équivalent du montage.
 - Établir l'équation différentielle régissant l'évolution de $i_H(t)$ et donner la solution générale de cette équation.
- L'interrupteur H est ouvert ($\alpha T < t < T$).
 - Quel est le sens du courant $i(t)$ juste après l'ouverture de l'interrupteur H ? Quel est le rôle de la diode D ? Donner le schéma équivalent au montage pour cet intervalle de temps.
 - Établir l'équation différentielle régissant l'évolution de $i_D(t)$ pour cet intervalle de temps et en déduire la solution générale.
- Représenter graphiquement l'évolution de la tension $u(t) = -u_D(t)$. Calculer sa valeur moyenne $\langle U \rangle$ en fonction de E et α .

- b) Calculer le rapport $\frac{\tau}{T}$ où $\tau = \frac{L}{R}$. On supposera que $\tau \ll T$. Donner une expression approchée des courants $i_H(t)$ et $i_D(t)$ en effectuant un développement limité au premier ordre.
- c) Donner l'allure des courants $I(t)$, $I_H(t)$ et $I_D(t)$. On notera $i_{\max}$ et $i_{\min}$ les valeurs maximales et minimales de $I(t)$.
- d) Montrer que la valeur moyenne de u_L est nulle. En déduire la valeur moyenne de i en fonction de α , E , E' et R .
- e) Expliquer le rôle de la bobine.

3. Quelques applications des photodiodes, d'après Centrale TSI 2006

Une photodiode est un composant électrooptique dont la caractéristique électrique dépend de la puissance lumineuse moyenne reçue au niveau de la surface sensible. L'équation de sa caractéristique en convention récepteur est donnée par $i(u) = I_0 \left(\exp\left(\frac{u}{V_0}\right) - 1 \right) - I_p$ en notant $I_0 = 10,0 \mu\text{A}$, $V_0 = 26,0 \text{ mV}$ et I_p le « photocourant » proportionnel à la puissance lumineuse reçue P_l avec un coefficient de proportionnalité $k = 0,50 \text{ A.W}^{-1}$.

1. Soit une photodiode recevant une puissance $P_l = 1,0 \text{ mW}$. Représenter graphiquement la caractéristique de la diode.
2. Exprimer la tension u_{C0} de la diode en circuit ouvert en fonction de V_0 , I_p et I_0 et donner sa valeur numérique.
3. Montrer qu'il existe un domaine du plan (u, i) où la photodiode fournit une puissance positive au circuit dans lequel elle se trouve.
4. Dans quelle région du plan, la photodiode peut-elle être considérée comme une source idéale de courant ?
5. Dans la suite, on adopte le modèle simplifié suivant : la caractéristique est composée de deux segments de droite : $i = -I_p - I_0$ pour $u < u_{C0}$ et $u = u_{C0}$ pour $i > -I_p - I_0$.
On connecte la photodiode à une résistance R_c . Quelle relation supplémentaire entre i et u a-t-on ?
6. Déterminer graphiquement les valeurs de l'intensité i et de la tension u . On distinguera deux cas en introduisant une valeur limite de la résistance R_0 qu'on exprimera en fonction de I_p , I_0 et u_{C0} .
7. Déterminer la puissance fournie par la photodiode en fonction de R_c , u_{C0} , I_0 et I_p .
8. Représenter graphiquement P en fonction de R_c .
9. Établir l'existence d'une valeur maximale de la puissance $P_{\max}$. On donnera la valeur R_{opt} de R_c pour laquelle elle est obtenue.
10. On définit le rendement de conversion de la photodiode par $\eta = \frac{P_{\max}}{P_l}$. Exprimer η en fonction de k , V_0 et $x = \frac{kP_l}{I_0}$.
11. Déterminer la valeur du rendement pour $P_l = 1,0 \text{ mW}$.

12. Étudier la limite x tendant vers l'infini. Que peut-on en déduire pour modèle utilisé pour décrire la caractéristique de la photodiode ?
13. Justifier que le modèle simplifié tend à surévaluer le rendement de conversion de la photodiode.
14. On associe maintenant en série N photodiodes recevant chacune la même puissance lumineuse P_l . Déterminer la tension de circuit ouvert u'_{C0} et le courant de court-circuit i'_{cc} du dipôle ainsi constitué.
15. Donner la puissance maximale $P'_{\max}$ que peut délivrer le générateur ainsi constitué en fonction de N et $P_{\max}$.
16. Déterminer en fonction de R_{opt} et N la résistance de charge optimale qui doit être connectée au générateur constitué des N photodiodes en série de manière à récupérer le maximum de puissance.
17. Reprendre les questions précédentes pour une association en parallèle des photodiodes. On notera u''_0 la tension de circuit ouvert, i''_{cc} la courant de court-circuit, $P''_{\max}$ la puissance maximale et R''_{opt} la résistance optimale.
18. On souhaite alimenter une résistance R_c de $1,0 \text{ k}\Omega$ à l'aide d'un ensemble de photodiodes associées en série ou en parallèle, chacune exposée à la même puissance lumineuse $P_l = 1,0 \text{ mW}$. Déterminer numériquement le nombre de photodiodes à employer ainsi que la façon optimale de les connecter (soit en série, soit en parallèle) pour recueillir le maximum de puissance dans la résistance. Donner alors le rendement en puissance de l'installation.
19. On utilise maintenant les photodiodes en détecteur pour mesurer P_l . On branche pour cela la photodiode en série avec une source idéale de tension $E = -15 \text{ V}$ et une résistance R . On admet que la photodiode est polarisée en inverse ($u < 0$) et qu'alors on peut la modéliser par l'association en parallèle d'une source idéale de courant $-I_p$, une résistance $R_d = 10,0 \text{ m}\Omega$ et une capacité $C_d = 2,0 \text{ pF}$.
- Pour le régime statique, on suppose que P_l est indépendante du temps. Déterminer la tension v_s aux bornes de R en fonction de E , I_p , R et R_d .
20. A quelle condition sur I_p , E et R_d peut-on considérer que v_s est proportionnelle à P_l ? On supposera que cette condition est vérifiée dans la suite et on notera $A = -\frac{v_s}{P_l}$. Exprimer A en fonction de k , R et R_d .
21. Déterminer en fonction de E et A la puissance lumineuse maximale $P_{l,\text{sat}}$ qu'on peut ainsi détecter tout en maintenant $u < 0$. Donner sa valeur numérique pour $R = 10 \text{ k}\Omega$.
22. Déterminer le seuil de détection σ du détecteur c'est-à-dire la puissance lumineuse donnant en sortie la plus petite tension considérée comme mesurable $|v_{SN}| = 1,0 \text{ mV}$. On donnera son expression en fonction de A et $|v_{SN}|$ ainsi que sa valeur numérique.
23. Déterminer l'équation différentielle entre $v_s(t)$ et $P_l(t)$ en régime dynamique. On fera apparaître une constante de temps $\tau = \frac{RR_dC_d}{R + R_d}$ et A .
24. La puissance lumineuse reçue s'écrit $P_l(t) = P_0 + \Delta P_l(t) = P_0 + \Delta P \cos \omega t$ en notant ΔP une constante. Justifier que $v_s(t) = v_{s0} + \Delta v_s(t)$ où v_{s0} est une constante à exprimer en fonction de A et P_0 et $\Delta v_s(t)$ une tension sinusoïdale de pulsation ω .

25. Exprimer la fonction de transfert $H = \frac{\Delta v_s}{\Delta P_l}$ en fonction de A , ω et τ .
26. Tracer, après l'avoir rapidement étudié, le diagramme de Bode et identifier la nature du filtrage effectué par la détection.
27. Déterminer l'expression et la valeur numérique de la pulsation de coupure ω_c .
28. On considère que le détecteur est exposé à une impulsion lumineuse telle que $P_l(t) = P_0$ pour $0 \leq t \leq \Delta t$ et $P_l(t) = 0$ sinon. Déterminer l'expression de $v_s(t)$ en supposant que $v_s(t < 0) = 0$ et représenter graphiquement $v_s(t)$.
29. On considère les deux instants t' et t'' avec $t' < t''$ pour lesquels $v_s(t) = -\frac{v_{s,\max}}{2}$ pour définir la largeur de l'impulsion détectée $\Delta t_d = t'' - t'$. Exprimer Δt_d en fonction de Δt et τ .
30. Étudier les cas $\Delta t \ll \tau$ et $\tau \ll \Delta t$.
31. Quelle est la plus petite valeur Δt_d qu'il est possible de détecter ? On pourra étudier les variations de Δt_d en fonction de Δt .
32. On utilise ce montage pour détecter de l'information transmise sous forme binaire avec des impulsions de largeur Δt . On appelle bande passante B le nombre maximal d'impulsions pouvant être détectées par unité de temps. Donner un ordre de grandeur de B en fonction de τ .
33. Montrer que le rapport $\frac{B}{\sigma}$ est une constante du détecteur et du niveau de bruit $|v_{SN}|$.
34. On veut détecter un signal avec une bande passante de 10^7 impulsions par seconde. Quelle doit être la puissance minimale disponible au niveau du détecteur ?

4. Étude d'une photopile solaire, d'après E3A MP 2005

L'une des applications du silicium de qualité électronique est la réalisation de cellules photovoltaïques ou photopiles solaires qui, associées sous forme de panneaux, permettent une transformation directe d'une énergie électromagnétique (rayonnement solaire) en énergie électrique de type continu directement utilisable. Ce capteur est constitué d'une jonction semi-conductrice de type P-N.

1. Caractéristique électrique d'une photopile

Une jonction P-N, polarisée par une tension V , possède une caractéristique courant-tension satisfaisant à l'équation :

$$I_d = I_s \left[\exp \left(\frac{qV}{kT} \right) - 1 \right]$$

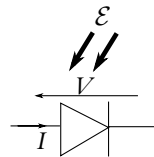


Figure 22.43

I_s étant le courant de saturation (ou courant d'obscurité des porteurs minoritaires), k la constante de Boltzmann et T la température. Si la surface supérieure du capteur d'aire Σ est soumise à un flux lumineux, on observe la création d'un courant électrique **inverse** (appelé

photocourant I_{ph}). La caractéristique de la photodiode éclairée est alors modélisée par l'expression : $I = I_d - I_{ph}$ avec $I_{ph} = \alpha \Sigma \mathcal{E}$ où $\mathcal{E}$ est la puissance lumineuse reçue par unité de surface et α un coefficient positif. La figure (22.44) montre cette caractéristique, sans éclairage, puis sous différents éclairages d'intensité croissante. Selon le quadrant de fonctionnement du dispositif, la photodiode est utilisée comme détecteur ou comme générateur de courant.

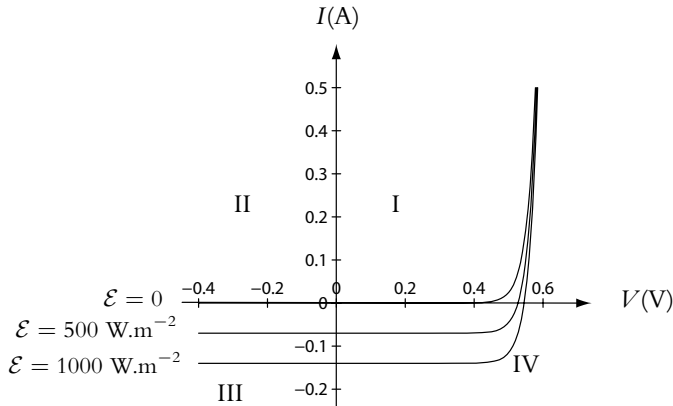


Figure 22.44

Données relatives au comportement électrique de la photopile :

$T = 300 \text{ K}$; $\mathcal{E} = 1 \text{ kW.m}^{-2}$; $I_s = 0,1 \text{ nA}$; $\Sigma = 4.10^{-4} \text{ m}^2$ et $\alpha = 0,35 \text{ A.W}^{-1}$.

a) Calculer le photocourant I_{ph} pour une puissance $\mathcal{E}$ de 1 kW/m^2 .

Préciser dans quel quadrant (I à IV, figure 22.44) la diode peut fonctionner en mode photovoltaïque. Retrouver, à l'aide de ce graphique, le résultat du calcul précédent.

b) Déterminer, en circuit ouvert, la tension de sortie notée V_{C0} ; préciser puis calculer sa valeur limite si $I_{ph} \gg I_s$; quelle puissance électrique le dispositif peut-il fournir à l'extérieur?

c) Si la photopile est en court-circuit, que vaut le courant de court-circuit noté I_{cc} ? Quelle est alors la puissance électrique fournie?

d) Une résistance de charge R_L étant placée aux bornes de la cellule, l'utilisation optimale du dispositif consiste à faire fonctionner cette charge sous une tension maximale V_M et une intensité maximale I_M .

Reproduire sur un schéma le tracé de la caractéristique de la photodiode (figure 22.44) pour une puissance $\mathcal{E}$ de 1 kW/m^2 , en prenant comme sens positif du courant celui du courant débité par la photopile. Déterminer et calculer V_M et I_M associés à une puissance électrique maximale $\mathcal{P}_M$.

Positionner les points associés à ces grandeurs, puis trouver la résistance de charge optimale R_{opt} .

e) Le rendement de conversion η de la cellule est défini comme le rapport de la puissance maximale pouvant être extraite, à la puissance du rayonnement incident sur la surface Σ de la cellule.

Écrire et calculer ce rendement; analyser le résultat obtenu.

2. Capteur solaire

Pour augmenter les performances de ce dispositif, on associe en série et en parallèle, respectivement n_s et n_p cellules rigoureusement identiques à celle étudiée.

a) Déterminer l'évolution de la caractéristique courant-tension résultant de l'association série, de l'association parallèle, puis de leur combinaison ? En déduire les valeurs V_{MC} et I_{Mc} qui en découlent pour $n_s = 36$ et $n_p = 6$.

Tracer la nouvelle caractéristique de fonctionnement de cette association.

b) Le capteur solaire étant couplé à un ventilateur (assimilé à une résistance de $22\ \Omega$), tracer sur ce dernier schéma la caractéristique de ce ventilateur et préciser le point de fonctionnement ; évaluer la puissance reçue par le ventilateur.

Reprendre la même étude dans le cas où la résistance du ventilateur ne vaut que $10\ \Omega$ et conclure qualitativement.

c) Que se serait-il passé si la photopile n'alimentait plus un récepteur mais une batterie ?

5. Étude d'un wattmètre électronique, d'après CCP PC 2003

1. Caractéristique d'une diode

Dans tout le problème, les amplificateurs qui vont être étudiés utilisent une diode, schématisée sur la figure (22.45), dont la caractéristique courant-tension a pour équation :

$$i = I_0 \left(\exp \left(\frac{v}{V_0} \right) - 1 \right)$$

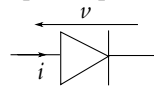


Figure 22.45

où i est l'intensité du courant traversant la diode, v la tension à ses bornes et I_0 et V_0 des constantes positives.

Pour les applications numériques, on prendra : $I_0 = 10\ \mu\text{A}$ et $V_0 = 25\ \text{mV}$.

Tracer qualitativement l'allure de la courbe $i(v)$.

2. Amplificateur logarithmique

On réalise la montage de la figure (22.46) :

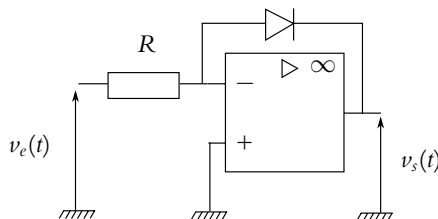


Figure 22.46

L'amplificateur opérationnel (A.O.) est supposé idéal et on note V_{sat} sa tension de saturation égale à $\pm 20\ \text{V}$; les sources de polarisation (continues) ne figurent pas sur les schémas.

a) L'amplificateur opérationnel étant supposé en régime linéaire, déterminer $v_s(t)$ en fonction de $v_e(t)$, V_0 , I_0 et R .

b) On suppose que $v_e(t) = V_e\sqrt{2}\sin(\omega t)$ et que $V_e\sqrt{2} > RI_0$ et on remarquera que $V_{\text{sat}} \gg V_0$.

Pour $0 < \omega t < \pi$, et compte-tenu des signes de v_e et de v_s , justifier s'il y aura ou non saturation de l'amplificateur.

Répondre à la même question pour $\pi < \omega t < 2\pi$.

c) Tracer l'allure des courbes $v_e(t)$ et $v_s(t)$ en fonction du temps, sur une période complète, pour $V_e\sqrt{2} > RI_0$.

3. Amplificateur opérationnel

On réalise la montage de la figure (22.47) :

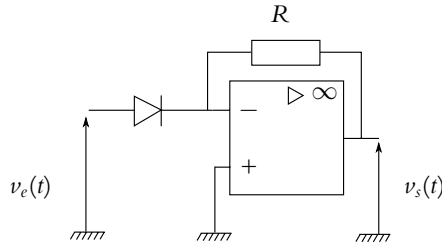


Figure 22.47

a) L'amplificateur opérationnel étant supposé idéal en régime linéaire, déterminer $v_s(t)$ en fonction de $v_e(t)$, V_0 , I_0 et R .

b) On suppose que $v_e(t) = V_e\sqrt{2}\sin(\omega t)$; pour $0 < \omega t < \pi$, donner la condition que $v_e(t)$ doit satisfaire pour que l'amplificateur soit effectivement en régime linéaire ; calculer la valeur numérique de la limite trouvée pour $v_e(t)$ à partir des valeurs suivantes : $R = 10^6 \Omega$ et $V_{\text{sat}} = 20 \text{ V}$.

c) Obtenir la condition entre V_{sat} et RI_0 nécessaire pour que l'amplificateur opérationnel soit en régime linéaire lorsque $\pi < \omega t < 2\pi$, quelle que soit la valeur de V_e .

4. Wattmètre électronique

On réalise le montage de figure (22.48) dans lequel les amplificateurs opérationnels sont en régime linéaire :

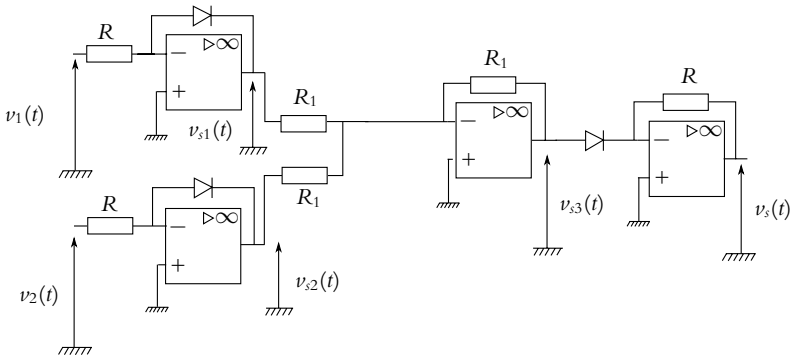


Figure 22.48

- a) Exprimer v_{s3} en fonction de v_{s1} et v_{s2} puis en fonction de v_1 et v_2 et des éléments du montage.
- b) En déduire la caractéristique de transfert $v_s = f(v_1, v_2)$ de ce montage, en fonction de I_0 et R .
- c) On considère que les tensions d'entrée sont de la forme :

$$v_1(t) = V_1 \sqrt{2} \cos(\omega t) \text{ et } v_2(t) = V_1 \sqrt{2} \cos(\omega t - \varphi)$$

Déterminer l'expression de la valeur moyenne dans le temps de la tension de sortie notée $\langle v_s \rangle$.

- d) Proposer un moyen pour mesurer $\langle v_s \rangle$.
- e) On considère un dipôle constitué de résistances, bobines et condensateurs, d'impédance complexe équivalente $\underline{Z}_{eq} = R_{eq} + jX_{eq}$ et alimenté par une tension $v(t) = V\sqrt{2} \cos(\omega t)$ (figure 22.49).

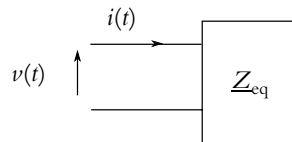


Figure 22.49

Exprimer la valeur moyenne P de la puissance instantanée reçue par le dipôle en fonction de V , R_{eq} et X_{eq} dite aussi « puissance active ».

- f) On veut mesurer cette puissance avec le montage de la figure (22.48) noté W ; on réalise dans ce but la montage de la figure (22.50).

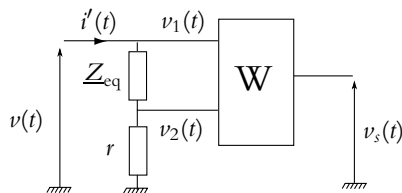


Figure 22.50

On considère que les intensités des deux entrées du wattmètre sont nulles.

Quelle est le rôle de la résistance r ? Comment doit-on choisir la valeur de celle-ci ?

Montrer que la puissance moyenne totale mesurée est de la forme : $P' = k \langle v_s \rangle$. Expliciter la constante k et exprimer P' en fonction de V , R_{eq} , r et X_{eq} .

Mécanique 2

Compléments de cinématique du point matériel

23

Au cours de la première période, seuls les mouvements plans ont été étudiés. Au cours de la deuxième période, on envisage des mouvements dans l'espace : les notions de cinématique du point matériel vont donc être rappelées et complétées pour l'espace.

1. Systèmes de coordonnées dans l'espace

1.1 Coordonnées cartésiennes

Les coordonnées cartésiennes déjà rencontrées dans le plan se généralisent à l'espace de la manière suivante.

a) Définition

La base de projection est définie par trois vecteurs unitaires formant une base directe. Le premier est choisi *a priori* et les autres s'en déduisent par la recherche du caractère direct de la base.

En notant $\vec{u}_x$, $\vec{u}_y$ et $\vec{u}_z$ les vecteurs unitaires sur chacun des axes de la base, on obtient :

$$\overrightarrow{OM} = x\vec{u}_x + y\vec{u}_y + z\vec{u}_z$$

Ce système de coordonnées est celui auquel on pense le plus souvent mais ce n'est pas forcément le plus adapté à la symétrie du problème, notamment lorsqu'elle est cylindrique ou sphérique.

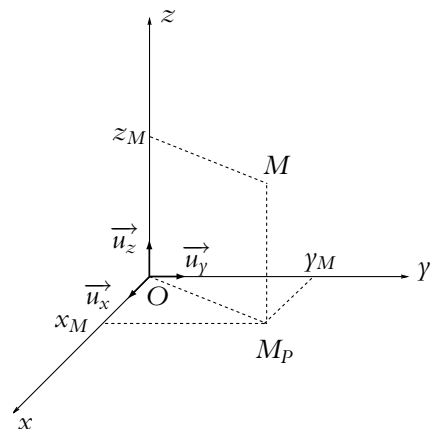


Figure 23.1 Coordonnées cartésiennes.

Pour éviter des confusions, sur les figures, les coordonnées cartésiennes de M seront en général notées x_M , y_M et z_M dans la suite du livre.

b) Déplacement élémentaire

On rappelle que le déplacement élémentaire s'obtient en faisant varier de manière élémentaire chacune des coordonnées du système de coordonnées considéré et sont particulièrement utiles pour la détermination de l'expression de la vitesse.

Ainsi, en coordonnées cartésiennes, le déplacement élémentaire d'un point M de coordonnées (x, y, z) correspond à son déplacement jusqu'au point M' $(x + dx, y + dy, z + dz)$. On a donc :

$$d\overrightarrow{OM} = \overrightarrow{OM'} - \overrightarrow{OM} = \overrightarrow{MM'} = dx\overrightarrow{u_x} + dy\overrightarrow{u_y} + dz\overrightarrow{u_z}$$

Ce résultat peut s'obtenir géométriquement en sommant les déplacements élémentaires liés à la variation d'une seule variable, les autres restant constantes. Ceci est possible car les variables sont indépendantes.

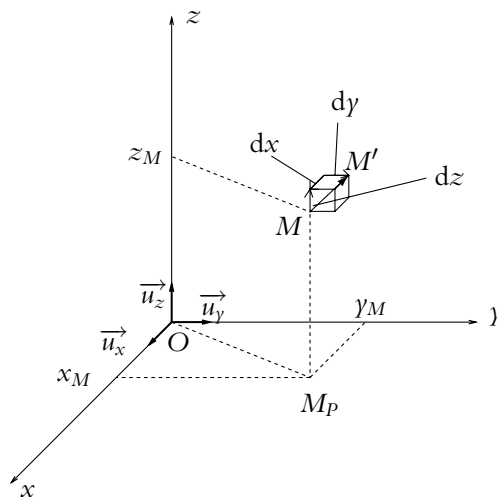


Figure 23.2 Déplacements élémentaires en coordonnées cartésiennes.

Ainsi, lorsque y et z (respectivement x et z ou x et y) restent fixes, le déplacement élémentaire se limite à $d\overrightarrow{OM} = dx\overrightarrow{u_x}$ (respectivement à $d\overrightarrow{OM} = dy\overrightarrow{u_y}$ ou $d\overrightarrow{OM} = dz\overrightarrow{u_z}$).

1.2 Coordonnées cylindriques

De même, on généralise les coordonnées polaires utilisées en première période par les coordonnées cylindriques en ajoutant le même axe Oz que pour les coordonnées cartésiennes. On obtient alors les résultats suivants.

a) Définition

On privilégie un axe (Oz) qui est souvent l'axe de symétrie cylindrique du problème. La position d'un point M est alors définie par l'altitude z et la distance r du point M à l'axe (Oz) .

On appelle *plan polaire* le plan perpendiculaire à l'axe (Oz) passant par M et on notera M_P le projeté orthogonal de M sur le plan (xOy) .

Dans le plan (xOy) , on repère M_P par ses coordonnées polaires r et θ . On rappelle que θ est un angle de rotation orienté appelé *angle polaire*.

Les coordonnées cylindriques de M sont donc (r, θ, z) . Pour éviter des confusions, elles seront souvent notées par la suite r_M , θ_M et z_M sur les figures.

On rappelle que r est une distance et est donc positif. On peut noter alors que pour décrire la totalité des points de l'espace (ou du plan), l'angle θ doit décrire un segment d'amplitude 2π . Habituellement, on prend θ dans l'intervalle $[0, 2\pi]$.

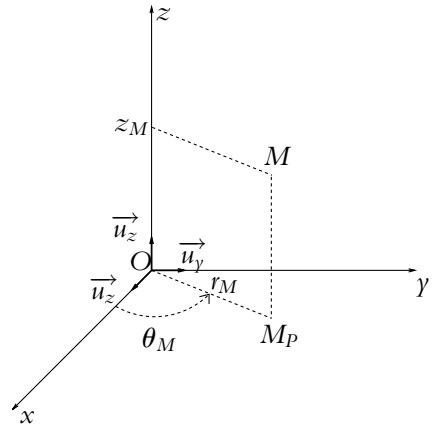


Figure 23.3 Coordonnées cartésiennes.

b) Lien avec les coordonnées cartésiennes

On peut relier les composantes en coordonnées cylindriques aux coordonnées cartésiennes. Il suffit pour cela de choisir l'axe (Ox) par exemple comme axe de référence pour l'angle de rotation θ et de projeter sur les axes (Ox) et (Oy) comme cela a déjà été fait pour établir le lien entre coordonnées polaires et cartésiennes :

$$\begin{cases} x = r \cos \theta \\ y = r \sin \theta \\ z = z \end{cases}$$

On peut inverser ces relations pour obtenir l'expression des coordonnées cylindriques en fonction des coordonnées cartésiennes :

$$\begin{cases} r = \sqrt{x^2 + y^2} \\ \theta = \arctan \frac{y}{x} [\pi] \\ z = z \end{cases}$$

La valeur de θ sera précisée dans l'intervalle $[0, 2\pi]$ en fonction des signes respectifs de x et y qui sont ceux de $\cos \theta$ et $\sin \theta$.

Ceci n'est possible qu'en dehors de l'origine O . En ce point, $x = y = 0$, donc la tangente de l'angle θ en O est une forme indéterminée.

c) Base locale cylindrique

On rappelle que les bases locales correspondent à des bases définies en chaque point M de l'espace et adaptées au type de coordonnées choisies. Ainsi pour définir la base locale des coordonnées cylindriques, on remplace les vecteurs de base $(\vec{u}_x, \vec{u}_y, \vec{u}_z)$ adaptés aux coordonnées cartésiennes par des vecteurs de base adaptés aux coordonnées cylindriques.

On rappelle également que le fait que la base « suive » le point M dans son mouvement définit une base mobile où les vecteurs de base bougent avec le point contrairement à la base des coordonnées cartésiennes qui est fixe.

On construit ainsi la base locale définie par les trois vecteurs suivants :

- le premier vecteur noté $\vec{u}_r$ est un vecteur unitaire dont la direction et le sens sont ceux du vecteur $\vec{OM}_P$,
- le second noté $\vec{u}_\theta$ est un vecteur unitaire dans le plan polaire obtenu par rotation autour de l'axe Oz de $+\frac{\pi}{2}$ à partir de $\vec{u}_r$,
- le troisième noté $\vec{u}_z$ (ou $\vec{k}$) est un vecteur unitaire suivant l'axe (Oz) .

Les vecteurs $\vec{u}_r$ et $\vec{u}_\theta$ dépendent du point M et se déplacent avec celui-ci, d'où la dénomination *locale* donnée à cette base.

Le vecteur-position $\vec{OM}$ s'écrit dans cette base : $\vec{OM} = r\vec{u}_r + z\vec{u}_z$.

On peut exprimer les vecteurs de cette base locale dans la base des coordonnées cartésiennes : il suffit de projeter les vecteurs $\vec{u}_r$ et $\vec{u}_\theta$ sur les axes (Ox) et (Oy) , l'axe (Oz) restant fixe.

On obtient :

$$\begin{cases} \vec{u}_r = \cos \theta \vec{u}_x + \sin \theta \vec{u}_y \\ \vec{u}_\theta = -\sin \theta \vec{u}_x + \cos \theta \vec{u}_y \end{cases}$$

Tout vecteur peut se décomposer sur cette base selon :

$$\vec{w} = w_r \vec{u}_r + w_\theta \vec{u}_\theta + w_z \vec{u}_z$$

w_r est appelé composante *radiale*, w_θ composante *orthoradiale* et w_z composante *axiale*.

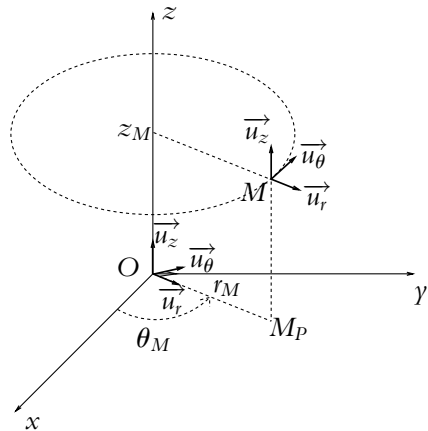


Figure 23.4 Base locale cylindrique.

d) Déplacement élémentaire

En coordonnées cylindriques, deux approches sont possibles pour déterminer le déplacement élémentaire.

1. Approche géométrique :

Les variables r, θ, z sont indépendantes : l'expression du déplacement élémentaire est donc la somme des déplacements élémentaires correspondant à la variation d'une seule variable.

- une variation dr de r correspond à un mouvement le long de la droite (OM) soit : $dr\vec{u}_r$,
- une variation $d\theta$ de θ correspond à un déplacement circulaire de centre O , de rayon r et d'angle $d\theta$ de longueur $rd\theta$ suivant $\vec{u}_\theta : rd\theta\vec{u}_\theta$,
- une variation dz de z correspond à un mouvement le long de l'axe (Mz) soit : $dz\vec{u}_z$.

Au final, le déplacement élémentaire s'écrit :

$$d\vec{OM} = dr\vec{u}_r + r d\theta\vec{u}_\theta + dz\vec{u}_z$$

2. Différenciation du vecteur $\vec{OM} : \vec{OM} = r\vec{u}_r + z\vec{u}_z$

$$\begin{aligned} d\vec{OM} &= d(r\vec{u}_r + z\vec{u}_z) \\ &= dr\vec{u}_r + r d\vec{u}_r + dz\vec{u}_z \\ &= dr\vec{u}_r + r d\theta\vec{u}_\theta + dz\vec{u}_z \end{aligned}$$

en utilisant le résultat établi en première période pour les coordonnées polaires :

$$d\vec{u}_r = d\theta\vec{u}_\theta$$

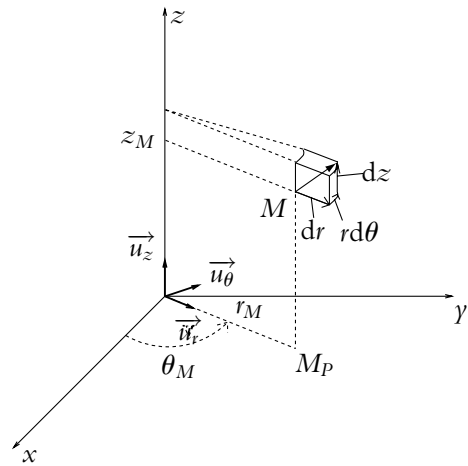


Figure 23.5 Déplacements élémentaires en coordonnées cylindriques.

1.3 Coordonnées sphériques

On peut dans l'espace définir un troisième système de coordonnées : les *coordonnées sphériques*. La base associée, comme pour les coordonnées cylindriques, est une base mobile.

a) Définition

La symétrie sphérique consiste à privilégier un point O et les rotations autour de ce point. La position est alors définie par la distance r du point M au point O et

deux angles θ et φ . L'angle θ est appelé *colatitude* et est défini comme l'angle entre un axe fixe noté (Oz) et la droite (OM) . L'angle φ est appelé *longitude* ou *azimut* et correspond à l'angle entre un axe (Ox) perpendiculaire à l'axe (Oz) précédent et la droite (OM_P) où M_P est la projection de M dans le plan perpendiculaire à (Oz) passant par O . Les coordonnées sphériques de M sont (r, θ, φ) , souvent notées r_M, θ_M et φ_M sur les figures.

Par définition, r est une distance et est donc positif. On peut noter que pour décrire la totalité des points de l'espace, les angles θ et φ ne peuvent pas appartenir tous les deux à un intervalle d'amplitude 2π : l'un des deux doit se limiter à une amplitude de π . Habituellement, on prend θ dans l'intervalle $[0, \pi]$ et φ dans l'intervalle $[0, 2\pi]$.

On peut également définir (Cf. figure 23.7) :

- le *parallèle* passant par M comme le cercle de centre H , projeté orthogonal de M sur l'axe Oz , passant par M et situé dans un plan parallèle à (xOy) ,
- le *méridien* passant par M comme le cercle de centre O de rayon r passant par M dans un plan perpendiculaire à (xOy) ,
- le *plan méridien* passant par M comme le plan contenant le méridien passant par M .

On notera que le plan méridien est un plan d'équation $\varphi = \varphi_0$ où φ_0 est une constante, r et θ correspondent alors aux coordonnées polaires du point M dans ce plan.

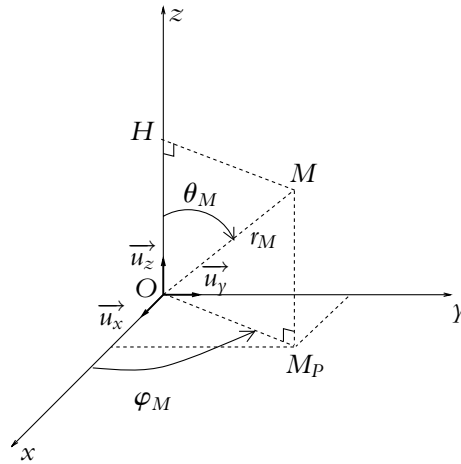


Figure 23.6 Coordonnées sphériques.

b) Lien avec les coordonnées cartésiennes

Comme pour les coordonnées cylindriques, on peut relier les composantes des coordonnées sphériques aux coordonnées cartésiennes. Il suffit pour cela de projeter le vecteur $\overrightarrow{OM}$ sur les axes (Ox) , (Oy) et (Oz) (tels que les vecteurs directeurs de ces axes $(\vec{u}_x, \vec{u}_y, \vec{u}_z)$ forment une base directe) :

$$\begin{cases} x = r \sin \theta \cos \varphi \\ y = r \sin \theta \sin \varphi \\ z = r \cos \theta \end{cases}$$

On peut inverser ces relations pour obtenir l'expression des coordonnées sphériques en fonction des coordonnées cartésiennes :

$$\begin{cases} r = \sqrt{x^2 + y^2 + z^2} \\ \theta = \arctan \frac{\sqrt{x^2 + y^2}}{z} [\pi] \\ \varphi = \arctan \frac{y}{x} [\pi] \end{cases}$$

θ et φ sont définis dans un intervalle de définition qui tient compte des signes respectifs de x , y et z qui sont ceux de $\cos \varphi$, $\sin \varphi$ et $\cos \theta$.

Ceci n'est pas possible en tout point :

- θ est indéterminé à l'origine où $x = y = z = 0$,
- φ est indéterminé sur l'axe (Oz) où $x = y = 0$.

Dans les deux cas, l'expression de la tangente des angles est alors une forme indéterminée.

c) Base locale sphérique

Il s'agit de la base locale définie par les trois vecteurs tels que :

- le premier vecteur noté $\vec{u}_r$ est un vecteur unitaire dont la direction et le sens sont ceux du vecteur $\vec{OM}$,
- le second noté $\vec{u}_\theta$ est un vecteur unitaire dans le plan méridien (Oz, OM) obtenu par rotation de $+\frac{\pi}{2}$ à partir de $\vec{u}_r$: il est donc tangent au méridien passant par M ,
- le troisième noté $\vec{u}_\varphi$ est un vecteur unitaire tel que $(\vec{u}_r, \vec{u}_\theta, \vec{u}_\varphi)$ soit une base directe : il est donc perpendiculaire au plan méridien et tangent au parallèle passant par M .

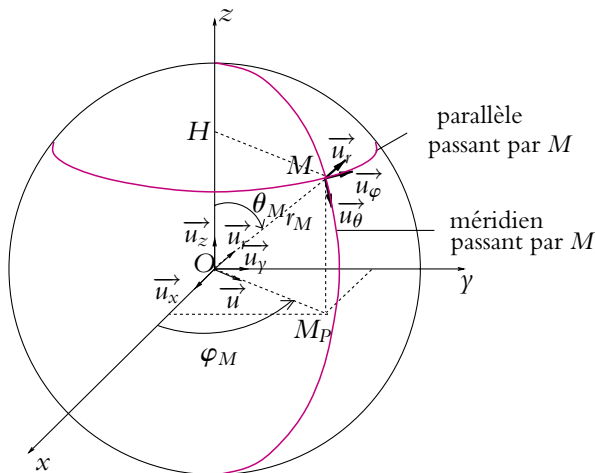


Figure 23.7 Base locale sphérique.

Les vecteurs $\vec{u}_r$, $\vec{u}_\theta$ et $\vec{u}_\varphi$ dépendent du point M : lorsque le point M se déplace, ces vecteurs changent également d'orientation, ce qui justifie la dénomination de base locale.

Le vecteur position $\vec{OM}$ s'écrit dans cette base : $\vec{OM} = r\vec{u}_r$.

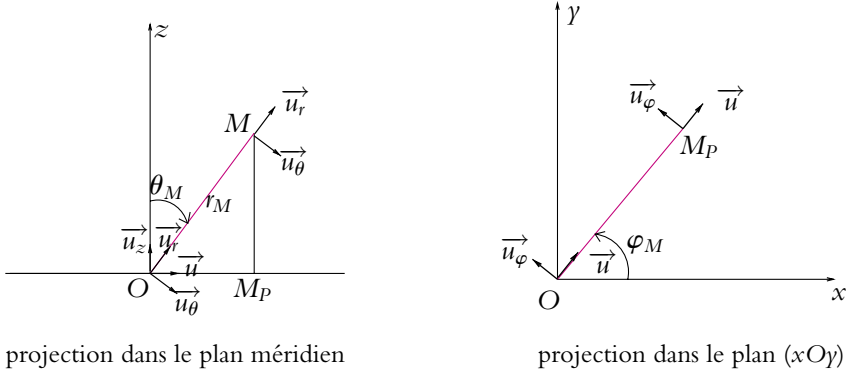


Figure 23.8 Projection de la base locale sphérique dans le plan méridien et dans le plan (xOy) .

On peut exprimer les vecteurs de cette base locale dans la base des coordonnées cartésiennes : il suffit de projeter les vecteurs $\vec{u}_r$, $\vec{u}_\theta$ et $\vec{u}_\varphi$ sur les axes (Ox) , (Oy) et (Oz) .

On a :

$$\begin{cases} \vec{u}_r = \sin \theta \vec{u} + \cos \theta \vec{u}_z = \sin \theta \cos \varphi \vec{u}_x + \sin \theta \sin \varphi \vec{u}_y + \cos \theta \vec{u}_z \\ \vec{u}_\theta = \cos \theta \vec{u} - \sin \theta \vec{u}_z = \cos \theta \cos \varphi \vec{u}_x + \cos \theta \sin \varphi \vec{u}_y - \sin \theta \vec{u}_z \\ \vec{u}_\varphi = -\sin \varphi \vec{u}_x + \cos \varphi \vec{u}_y \end{cases} \quad (23.1)$$

en notant :

$$\vec{u} = \cos \varphi \vec{u}_x + \sin \varphi \vec{u}_y$$

Tout vecteur peut se décomposer sur cette base locale selon :

$$\vec{w} = w_r \vec{u}_r + w_\theta \vec{u}_\theta + w_\varphi \vec{u}_\varphi$$

w_r est appelé composante *radiale*, w_θ et w_φ n'ont pas de nom consacré.

On peut cependant remarquer que si M est un point à la surface de la Terre, w_θ est la composante nord-sud et w_φ la composante est-ouest.

d) Déplacement élémentaire

En coordonnées sphériques, on obtient l'expression du déplacement élémentaire par l'une des deux approches précédemment développées en coordonnées cylindriques.

1. Approche géométrique :

Les variables r, θ, φ sont indépendantes : l'expression du déplacement élémentaire est donc la somme des déplacements élémentaires correspondant à la variation d'une seule variable.

- une variation dr de r correspond à un mouvement le long de la droite (OM) donc de $dr \vec{u}_r$,
- une variation $d\theta$ de θ correspond à un déplacement circulaire de centre O , de rayon r et d'angle $d\theta$ de longueur $rd\theta$ suivant $\vec{u}_\theta$. Ceci se passe dans le plan défini par les droites (OM) et (Oz) ,
- une variation $d\varphi$ de φ correspond à un mouvement circulaire perpendiculairement à l'axe (Oz) autour de cet axe. Le rayon du cercle est $r \sin \theta$ et l'angle de rotation $d\varphi$. La longueur du déplacement est donc $r \sin \theta d\varphi$ suivant $\vec{u}_\varphi$.

Au final, le déplacement élémentaire s'écrit :

$$d\vec{OM} = dr \vec{u}_r + r d\theta \vec{u}_\theta + r \sin \theta d\varphi \vec{u}_\varphi$$

2. Différenciation du vecteur

$\vec{OM} = r \vec{u}_r$: il s'agit d'une méthode délicate qu'on évitera autant que possible.

$$\begin{aligned} d\vec{OM} &= d(r \vec{u}_r) \\ &= dr \vec{u}_r + r d\vec{u}_r \\ &= dr \vec{u}_r + r d\theta \vec{u}_\theta + r \sin \theta d\varphi \vec{u}_\varphi \end{aligned}$$

En effet, on a :

$$d\vec{u}_r = \frac{\partial \vec{u}_r}{\partial \theta} d\theta + \frac{\partial \vec{u}_r}{\partial \varphi} d\varphi$$

Or :

$$\vec{u}_r = \sin \theta \cos \varphi \vec{u}_x + \sin \theta \sin \varphi \vec{u}_y + \cos \theta \vec{u}_z$$

d'où on peut déduire :

$$\frac{\partial \vec{u}_r}{\partial \theta} = \cos \theta \cos \varphi \vec{u}_x + \cos \theta \sin \varphi \vec{u}_y - \sin \theta \vec{u}_z = \vec{u}_\theta$$

$$\frac{\partial \vec{u}_r}{\partial \varphi} = -\sin \theta \sin \varphi \vec{u}_x + \sin \theta \cos \varphi \vec{u}_y = \sin \theta \vec{u}_\varphi$$

et

$$d\vec{u}_r = d\theta \vec{u}_\theta + \sin \theta d\varphi \vec{u}_\varphi$$

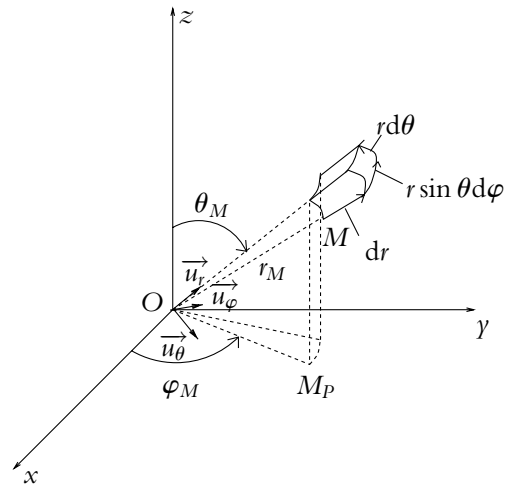


Figure 23.9 Déplacements élémentaires en coordonnées sphériques.

2. Dérivée d'un vecteur unitaire tournant par rapport à l'angle de rotation

Pour les bases locales en coordonnées cylindriques ou sphériques, les vecteurs $\vec{u}_r$, $\vec{u}_\theta$ et $\vec{u}_\varphi$ changent lorsque le point M se déplace. On s'intéresse ici aux modifications de ces vecteurs lors d'un déplacement du point M . On généralise ainsi ce qui a été fait pour les coordonnées polaires.

2.1 Cas des coordonnées cylindriques

Quand $M(r, \theta, z)$ devient $M'(r + dr, \theta + d\theta, z + dz)$, les vecteurs $\vec{u}_r$ et $\vec{u}_\theta$ deviennent respectivement $\vec{u}_r + d\vec{u}_r$ et $\vec{u}_\theta + d\vec{u}_\theta$.

Or les expressions de $\vec{u}_r$ et $\vec{u}_\theta$ établies dans la base fixe $(\vec{u}_x, \vec{u}_y, \vec{u}_z)$ ne dépendent que de θ :

$$\begin{cases} \vec{u}_r = \cos \theta \vec{u}_x + \sin \theta \vec{u}_y \\ \vec{u}_\theta = -\sin \theta \vec{u}_x + \cos \theta \vec{u}_y \end{cases}$$

Les variations induites par le déplacement de M ne sont donc fonction que de θ . En effet, quand r varie de dr à θ constant, la direction du vecteur $\overrightarrow{OM}_P$ ne change pas, la base locale non plus.

En dérivant par rapport à θ les expressions précédentes, on déduit :

$$\begin{cases} \frac{d\vec{u}_r}{d\theta} = \vec{u}_\theta \\ \frac{d\vec{u}_\theta}{d\theta} = -\vec{u}_r \end{cases} \quad \text{puis} \quad \begin{cases} d\vec{u}_r = d\theta \vec{u}_\theta \\ d\vec{u}_\theta = -d\theta \vec{u}_r \end{cases}$$

2.2 Cas des coordonnées sphériques

Quand $M(r, \theta, \varphi)$ devient $M'(r + dr, \theta + d\theta, \varphi + d\varphi)$, les vecteurs $\vec{u}_r$, $\vec{u}_\theta$ et $\vec{u}_\varphi$ deviennent respectivement $\vec{u}_r + d\vec{u}_r$, $\vec{u}_\theta + d\vec{u}_\theta$ et $\vec{u}_\varphi + d\vec{u}_\varphi$. Or les expressions de $\vec{u}_r$, $\vec{u}_\theta$ et $\vec{u}_\varphi$ établies dans la base fixe $(\vec{u}_x, \vec{u}_y, \vec{u}_z)$ ne dépendent que de θ et de φ : les variations induites par le déplacement de M ne sont donc fonction que de ces deux variables.

Des expressions de $\vec{u}_r$, $\vec{u}_\theta$ et $\vec{u}_\varphi$, on déduit :

$$\begin{cases} \left(\frac{\partial \vec{u}_r}{\partial \theta} \right)_{r, \varphi} = \vec{u}_\theta \\ \left(\frac{\partial \vec{u}_\theta}{\partial \theta} \right)_{r, \varphi} = -\vec{u}_r \\ \left(\frac{\partial \vec{u}_r}{\partial \varphi} \right)_{r, \theta} = \vec{u}_\varphi \\ \left(\frac{\partial \vec{u}_\varphi}{\partial \varphi} \right)_{r, \theta} = -\vec{u} = -\cos \varphi \vec{u}_x - \sin \varphi \vec{u}_y = -\sin \theta \vec{u}_r - \cos \theta \vec{u}_\theta \end{cases}$$

2.3 Cas général d'un vecteur unitaire

De ce qui précède, on remarque qu'on obtient le résultat suivant :

La dérivée d'un vecteur unitaire tournant par rapport à l'angle de rotation est un vecteur unitaire qui lui est directement orthogonal, dans le plan de rotation.

On peut établir qu'il s'agit d'un résultat tout à fait général. Soit $\vec{u}$ un vecteur unitaire. On a $\vec{u}^2 = 1$, ce qui implique en dérivant par rapport à l'angle de rotation α : $\vec{u} \cdot \frac{d\vec{u}}{d\alpha} = 0$. Les vecteurs $\vec{u}$ et $\frac{d\vec{u}}{d\alpha}$ sont donc orthogonaux.

Le fait que $\vec{u}$ et $\frac{d\vec{u}}{d\alpha}$ forment une base directe est lié à l'orientation de α dans le sens direct choisi.

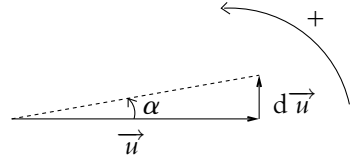


Figure 23.10 Dérivée d'un vecteur tournant par rapport à l'angle de rotation.

3. Expressions de la vitesse et de l'accélération

3.1 Coordonnées cartésiennes

On part du vecteur-position :

$$\overrightarrow{OM} = x(t)\vec{u}_x + y(t)\vec{u}_y + z(t)\vec{u}_z$$

et on dérive par rapport au temps sachant que les vecteurs $\vec{u}_x$, $\vec{u}_y$ et $\vec{u}_z$ sont constants au cours du temps.

Donc le vecteur-vitesse s'écrit :

$$\vec{v} = \frac{d\overrightarrow{OM}}{dt} = \dot{x}\vec{u}_x + \dot{y}\vec{u}_y + \dot{z}\vec{u}_z$$

en notant $\dot{x} = \frac{dx}{dt}$, $\dot{y} = \frac{dy}{dt}$ et $\dot{z} = \frac{dz}{dt}$.

De même pour l'accélération :

$$\vec{a} = \frac{d\vec{v}}{dt} = \ddot{x}\vec{u}_x + \ddot{y}\vec{u}_y + \ddot{z}\vec{u}_z$$

en notant $\ddot{x} = \frac{d^2x}{dt^2}$, $\ddot{y} = \frac{d^2y}{dt^2}$ et $\ddot{z} = \frac{d^2z}{dt^2}$.

3.2 Coordonnées cylindriques

Le vecteur-position s'écrit :

$$\overrightarrow{OM} = r(t)\vec{u}_r(t) + z(t)\vec{u}_z$$

On en déduit l'expression du vecteur-vitesse :

$$\begin{aligned}
 \vec{v} &= \frac{d\vec{OM}}{dt} \\
 &= \dot{r}\vec{u}_r + r\frac{d\vec{u}_r}{dt} + \dot{z}\vec{u}_z \\
 &= \dot{r}\vec{u}_r + r\frac{d\vec{u}_r}{d\theta}\frac{d\theta}{dt} + \dot{z}\vec{u}_z \\
 &= \dot{r}\vec{u}_r + r\dot{\theta}\vec{u}_\theta + \dot{z}\vec{u}_z
 \end{aligned}$$

en utilisant les résultats sur les dérivées par rapport à un angle d'un vecteur unitaire.

On peut retrouver ces résultats géométriquement : il suffit de sommer les déplacements élémentaires correspondant aux variations de chacune des variables pendant un intervalle de temps dt . On se reportera aux résultats obtenus pour les expressions des déplacements élémentaires en coordonnées cylindriques au paragraphe 1.2.d) de ce chapitre.

Pour l'accélération, on dérive l'expression de la vitesse :

$$\begin{aligned}
 \vec{a} &= \frac{d\vec{v}}{dt} \\
 &= \ddot{r}\vec{u}_r + \dot{r}\frac{d\vec{u}_r}{dt} + \dot{r}\dot{\theta}\vec{u}_\theta + r\frac{d\dot{\theta}}{dt}\vec{u}_\theta + r\dot{\theta}\frac{d\vec{u}_\theta}{dt} + \ddot{z}\vec{u}_z \\
 &= \ddot{r}\vec{u}_r + \dot{r}\frac{d\vec{u}_r}{d\theta}\frac{d\theta}{dt} + \dot{r}\dot{\theta}\vec{u}_\theta + r\ddot{\theta}\vec{u}_\theta + r\dot{\theta}\frac{d\vec{u}_\theta}{d\theta}\frac{d\theta}{dt} + \ddot{z}\vec{u}_z \\
 &= \ddot{r}\vec{u}_r + \dot{r}\vec{u}_\theta\dot{\theta} + \dot{r}\dot{\theta}\vec{u}_\theta + r\ddot{\theta}\vec{u}_\theta + r\dot{\theta}(-\vec{u}_r)\dot{\theta} + \ddot{z}\vec{u}_z \\
 &= (\ddot{r} - r\dot{\theta}^2)\vec{u}_r + (2\dot{r}\dot{\theta} + r\ddot{\theta})\vec{u}_\theta + \ddot{z}\vec{u}_z
 \end{aligned}$$

3.3 Coordonnées sphériques

Le vecteur-position s'écrit :

$$\vec{OM} = r(t)\vec{u}_r(t)$$

On en déduit l'expression du vecteur-vitesse :

$$\begin{aligned}
 \vec{v} &= \frac{d\vec{OM}}{dt} \\
 &= \dot{r}\vec{u}_r + r\frac{d\vec{u}_r}{dt} \\
 &= \dot{r}\vec{u}_r + r\frac{\partial\vec{u}_r}{\partial\theta}\frac{\partial\theta}{\partial t} + r\frac{d\vec{u}_r}{d\varphi}\frac{d\varphi}{dt} \\
 &= \dot{r}\vec{u}_r + r\dot{\theta}\vec{u}_\theta + r\sin\theta\dot{\varphi}\vec{u}_\varphi
 \end{aligned}$$

en utilisant les résultats sur les dérivées par rapport à un angle d'un vecteur unitaire.

On peut retrouver ces résultats géométriquement : il suffit de sommer les déplacements élémentaires correspondant aux variations de chacune des variables pendant un intervalle de temps dt . On se reportera aux résultats obtenus pour les expressions des déplacements élémentaires en coordonnées sphériques au paragraphe 1.3.d) de ce chapitre.

L'expression de l'accélération en coordonnées sphériques n'est que très rarement employée. On peut l'obtenir en dérivant l'expression obtenue pour la vitesse ; son expression¹ est :

$$\begin{aligned}\vec{a} &= \frac{d\vec{v}}{dt} \\ &= (\ddot{r} - r\dot{\theta}^2 - r\sin^2\theta\dot{\varphi}^2)\vec{u}_r + (2\dot{r}\dot{\theta} + r\ddot{\theta} - r\sin\theta\cos\theta\dot{\varphi}^2)\vec{u}_\theta \\ &\quad + (2\dot{r}\sin\theta\dot{\varphi} + r\sin\theta\ddot{\varphi} + 2r\cos\theta\dot{\theta}\dot{\varphi})\vec{u}_\varphi\end{aligned}$$

4. Compléments sur le mouvement circulaire

4.1 Rappels

On revient sur le mouvement circulaire pour donner les expressions de la vitesse et de l'accélération en introduisant un vecteur rotation et en utilisant la notion de produit vectoriel.

Du fait de la symétrie du problème, les coordonnées les plus adaptées sont les coordonnées cylindriques d'axe (Oz) qui généralisent les coordonnées polaires utilisées en première période. On rappelle que, dans cette base de projection, les équations horaires du mouvement sont :

$$\begin{cases} r = R = \text{constante} \\ \theta(t) = (Ox, \overrightarrow{OM}) \end{cases}$$

et que le vecteur vitesse s'écrit :

$$\vec{v} = r\dot{\theta}\vec{u}_\theta$$

4.2 Vecteur rotation

Il est souvent commode d'introduire un *vecteur rotation* $\vec{\omega}$ dont le module est la vitesse angulaire et dont la direction est celle de l'axe (Oz). Le sens du vecteur rotation est celui de $\vec{u}_z$ avec l'orientation choisie pour θ :

$$\vec{\omega} = \omega\vec{u}_z = \dot{\theta}\vec{u}_z$$

¹ On laisse le soin au lecteur de faire le calcul à titre d'entraînement.

On notera que ω est une grandeur algébrique : tous les sens de rotation possibles sont donc décrits par cette définition.

Il est alors possible d'exprimer le vecteur vitesse en fonction du vecteur position et du vecteur rotation. On se place dans le cas d'un sens direct de parcours.

$$\overrightarrow{OM} = r \vec{u}_r \quad \text{et} \quad \vec{\omega} = \omega \vec{u}_z$$

$$\vec{\omega} \wedge \overrightarrow{OM} = r \omega \vec{u}_z \wedge \vec{u}_r = r \omega \vec{u}_\theta$$

Donc

$$\vec{v} = \vec{\omega} \wedge \overrightarrow{OM}$$

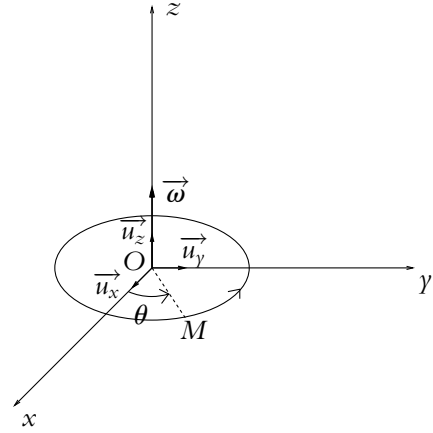


Figure 23.11 Vecteur rotation.

4.3 Expression de l'accélération

En première période, on avait obtenu l'expression de l'accélération en dérivant l'expression de la vitesse en coordonnées cylindriques par rapport au temps :

$$\vec{a} = \frac{d\vec{v}}{dt} = r\ddot{\theta}\vec{u}_\theta - r\dot{\theta}^2\vec{u}_r$$

Une autre méthode consiste à dériver l'expression de la vitesse qui vient d'être obtenue en fonction du vecteur rotation :

$$\vec{a} = \frac{d}{dt} (\vec{\omega} \wedge \overrightarrow{OM}) = \frac{d\vec{\omega}}{dt} \wedge \overrightarrow{OM} + \vec{\omega} \wedge \frac{d\overrightarrow{OM}}{dt} = \frac{d\vec{\omega}}{dt} \wedge \overrightarrow{OM} + \vec{\omega} \wedge \vec{v}$$

soit, en tenant compte du fait que $\vec{v} = r\omega\vec{u}_\theta$ pour un mouvement circulaire et en utilisant les coordonnées cylindriques :

$$\vec{a} = r \frac{d\omega}{dt} \vec{u}_z \wedge \vec{u}_r + r\omega^2 \vec{u}_z \wedge \vec{u}_\theta = -r\omega^2 \vec{u}_r + r \frac{d\omega}{dt} \vec{u}_\theta$$

On peut noter que dans le cas d'un mouvement circulaire uniforme, on a : $\frac{d\omega}{dt} = 0$ soit $\omega = \dot{\theta}$ constant. Par conséquent, l'accélération est uniquement radiale.

A. Applications directes du cours

1. Mouvement hélicoïdal

Un point matériel se déplace le long d'une hélice circulaire. Son mouvement est donné en coordonnées cylindriques par :

$$\begin{cases} r = R \\ \theta = \omega t \\ z = ht \end{cases}$$

où R , h et ω sont des constantes.

1. Donner l'expression de la vitesse.
2. En déduire que le module de la vitesse est constant.
3. Exprimer l'accélération.

B. Exercices et problèmes

1. Trajectoire d'un bateau

On s'intéresse à une planète semblable à la Terre mais entièrement recouverte d'eau. On assimile la planète à une sphère de rayon R . Un bateau se déplace à la surface de cette planète, avec une vitesse constante V_0 et toujours en direction du Nord-Est.

1. Exprimer les différentes composantes de la vitesse en projection dans la base des coordonnées sphériques.
2. Déterminer par intégration des équations précédentes, les expressions de la colatitude θ et de la longitude φ en fonction du temps.
3. Quelle est la date d'arrivée au pôle ?

24

Oscillations mécaniques forcées

Ce chapitre complète l'étude faite en première période sur les oscillations mécaniques harmoniques et amorties par frottement visqueux. On s'intéresse à l'influence d'une excitation harmonique sur un oscillateur amorti. En effet, la présence d'effets dissipatifs conduit à un amortissement puis à une disparition des oscillations. Si on souhaite observer des oscillations, il est donc nécessaire de les entretenir à l'aide d'une excitation extérieure : on obtient des oscillations forcées.

1. Position du problème et équation du mouvement

1.1 Entretien des oscillations

On a vu que l'amortissement des oscillations est lié à l'existence d'une force de frottement et à une diminution de l'énergie mécanique sous forme de chaleur dissipée lors du travail de la force de frottement.

Pour conserver les oscillations, il faut compenser cette perte en fournissant de l'énergie *via* une force extérieure.

1.2 Équation du mouvement

- Système : point matériel de masse m ,
- Référentiel du laboratoire considéré comme galiléen,
- Bilan des forces :
 - force de rappel $-kx\vec{u}_x$,
 - force de frottement $-\lambda\dot{x}\vec{u}_x$,
 - poids qui est compensé soit par la réaction de l'axe dans le cas où le pendule est horizontal, soit « par la position d'équilibre » dans le cas où le pendule est

vertical (Cf. exemple du pendule élastique vertical dans le paragraphe consacré à l'oscillateur harmonique en première période),

○ force extérieure, souvent dite excitatrice, $\vec{f}(t) = f(t)\vec{u}_x$.

- Principe fondamental de la dynamique projeté sur $\vec{u}_x$:

$$m\ddot{x} = -kx - \lambda\dot{x} + f(t)$$

L'équation différentielle du mouvement est donc

$$\ddot{x} + \frac{\lambda}{m}\dot{x} + \frac{k}{m}x = \ddot{x} + 2\xi\dot{x} + \omega_0^2x = \frac{f(t)}{m}$$

en posant $\omega_0 = \sqrt{\frac{k}{m}}$ et $\xi = \frac{\lambda}{2m}$.

On étudie donc la même équation différentielle que dans le cas des oscillations amorties avec un second membre correspondant à la force excitatrice qui dépend du temps.

1.3 Cas d'une excitation sinusoïdale

Dans la suite, on se limitera aux excitations sinusoïdales

$$f(t) = F \cos \omega t$$

Outre le fait que ce type d'excitation est important pour lui-même (vibrations d'une machine tournante, mouvement d'un électron dans un champ magnétique...), cette étude revêt un intérêt théorique capital. En effet, comme on l'a déjà vu en électricité, la fonction $f(t)$ peut s'écrire comme une superposition de fonctions sinusoïdales (discrète ou continue selon que $f(t)$ est périodique ou non). Le fait que l'équation différentielle soit linéaire, autrement dit que le principe de superposition puisse s'appliquer, permet d'écrire la solution pour $f(t)$ comme la somme des solutions obtenues pour chaque terme de la décomposition. Il est alors nécessaire de déterminer la réponse à chaque terme de la décomposition, à savoir à une excitation sinusoïdale. Ceci donne une importance considérable à l'étude de l'excitation sinusoïdale.

2. Étude de la solution

2.1 Régimes libre, transitoire et permanent

La solution d'une équation différentielle du type :

$$m\ddot{x} + \lambda\dot{x} + kx = F \cos \omega t$$

s'écrit sous la forme :

$$x(t) = x_H(t) + x_P(t)$$

où

- $x_H(t)$ correspond à la solution générale de l'équation homogène associée,
- $x_P(t)$ est une solution particulière de l'équation complète.

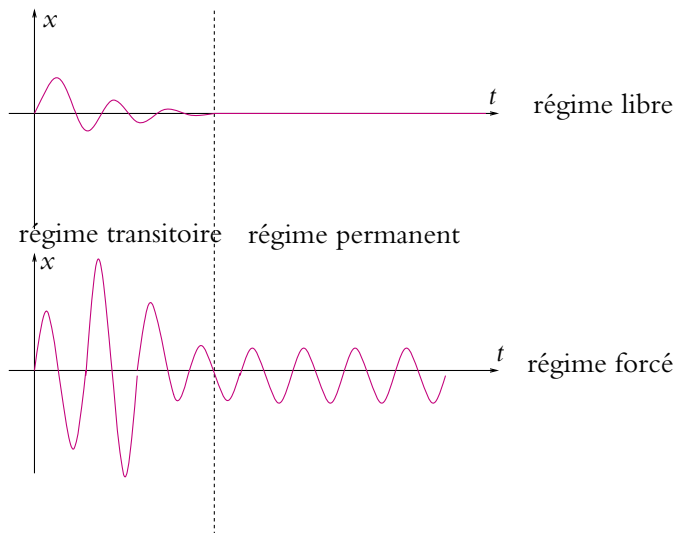


Figure 24.1 Régimes transitoire et permanent d'un régime libre et d'un régime forcé.

$x_H(t)$ est la solution obtenue en première période lors de l'étude des oscillations amorties. Il s'agit donc du régime libre. Dans tous les cas (pseudo-périodique, apériodique, critique), la solution tend vers 0 pour un temps infini, à savoir grand devant le temps de relaxation défini par $\tau = \frac{m}{\lambda}$. On notera que plus l'amortissement est grand, plus la force de frottement est importante et plus la durée caractéristique est courte. Cette partie de la solution a une durée de vie finie et devient négligeable au-delà d'un certain temps. Cette durée correspond au régime transitoire que l'on peut définir comme la partie de la solution qui apparaît lors d'un changement de régime et s'amortit au cours du temps.

Lorsque le régime transitoire a disparu, la solution s'identifie à $x_P(t)$ qui est la solution imposée par l'excitation $f(t)$. On dit qu'on a un régime forcé. On peut établir que lorsque $f(t)$ est sinusoïdale, la solution $x_P(t)$ l'est aussi et que la fréquence de la solution est la même que celle de l'excitation.

2.2 Détermination de la solution forcée

On cherche une solution particulière de l'équation globale sous la forme :

$$x_P(t) = X \cos(\omega t + \varphi)$$

On passe en notation complexe pour déterminer l'amplitude X et la phase φ . La représentation complexe de $x_P(t)$ s'écrit :

$$\underline{x} = X \exp(j(\omega t + \varphi)) = \underline{X} \exp(j\omega t) \quad \text{où} \quad \underline{X} = X \exp(j\varphi)$$

alors $x_P(t)$ est la partie réelle de $\underline{x}$. De même, l'excitation vérifie :

$$F \cos \omega t = \operatorname{Re}(F \exp(j\omega t))$$

L'intérêt de la notation complexe est le même qu'en électricité : les dérivées temporelles se traduisent en notation complexe par une multiplication par $j\omega$. Cela permet de transformer l'équation différentielle en équation algébrique :

$$m\underline{\ddot{x}} + \lambda\underline{\dot{x}} + k\underline{x} = (-m\omega^2 + j\omega\lambda + k) \underline{X} \exp(j\omega t) = F \exp(j\omega t)$$

Soit en simplifiant par $\exp(j\omega t)$ qui ne peut s'annuler :

$$(-m\omega^2 + j\omega\lambda + k) \underline{X} = F$$

L'amplitude complexe est donc :

$$\underline{X} = X \exp(j\varphi) = \frac{F}{k - m\omega^2 + j\omega\lambda} = \frac{F}{m} \frac{1}{\omega_0^2 - \omega^2 + 2j\xi\omega}$$

ce qui permet d'en déduire les expressions de l'amplitude et de la phase en fonction de la pulsation :

$$\left\{ \begin{array}{l} X = \frac{\frac{F}{m}}{\sqrt{(\omega_0^2 - \omega^2)^2 + 4\xi^2\omega^2}} \\ \tan \varphi = -\frac{2\xi\omega}{\omega_0^2 - \omega^2} \quad \text{et} \quad \sin \varphi < 0 \end{array} \right.$$

Si on utilise d'autres paramètres que ξ et ω_0 , on obtient :

1. en fonction de $\tau = \frac{m}{\lambda} = \frac{1}{2\xi}$ et de ω_0^2 :

$$\underline{X} = \frac{F}{m} \frac{1}{\omega_0^2 - \omega^2 + j\frac{\omega}{\tau}}$$

d'où :

$$\left\{ \begin{array}{l} X = \frac{\frac{F}{m}}{\sqrt{(\omega_0^2 - \omega^2)^2 + \frac{\omega^2}{\tau^2}}} \\ \tan \varphi = -\frac{\omega}{\tau(\omega_0^2 - \omega^2)} \end{array} \right. \quad \text{et} \quad \sin \varphi < 0$$

2. en fonction de $Q = \omega_0 \tau = \frac{\omega_0}{2\xi}$ et de ω_0^2 :

$$\underline{X} = \frac{F}{m} \frac{1}{\omega_0^2 - \omega^2 + j \frac{\omega \omega_0}{Q}}$$

d'où :

$$\left\{ \begin{array}{l} X = \frac{\frac{F}{m}}{\sqrt{(\omega_0^2 - \omega^2)^2 + \frac{\omega_0^2 \omega^2}{Q^2}}} \\ \tan \varphi = -\frac{1}{Q \left(\frac{\omega_0}{\omega} - \frac{\omega}{\omega_0} \right)} \end{array} \right. \quad \text{et} \quad \sin \varphi < 0$$

3. Étude de l'amplitude en fonction de la fréquence. Résonance en amplitude

On étudie dans ce paragraphe l'amplitude en fonction de la fréquence ν ou de la pulsation ω ($\omega = 2\pi\nu$).

3.1 Notion de résonance

On cherche le maximum de l'amplitude en fonction de la fréquence lorsqu'il existe. L'amplitude X est maximale lorsque la quantité $A(\omega) = (\omega_0^2 - \omega^2)^2 + 4\xi^2\omega^2$ est minimale.

$$\frac{dA}{d\omega} = 4\omega(-\omega_0^2 + \omega^2 + 2\xi^2)$$

Cette quantité s'annule pour $\omega = 0$ ou $\omega^2 = \omega_0^2 - 2\xi^2$. La deuxième solution ne donnera des valeurs réelles que si

$$\omega_0^2 > 2\xi^2 \quad \text{soit} \quad \xi\sqrt{2} < \omega_0$$

Cette condition $\xi\sqrt{2} < \omega_0$ est équivalente à $Q = \omega_0 \tau > \frac{1}{\sqrt{2}}$.

- si $\xi\sqrt{2} < \omega_0$ ou $Q = \omega_0\tau > \frac{1}{\sqrt{2}}$, alors on a deux solutions possibles :

$$\omega = 0 \quad \text{et} \quad \omega = \omega_m = \sqrt{\omega_0^2 - 2\xi^2} = \sqrt{\omega_0^2 - \frac{1}{2\tau^2}} = \omega_0 \sqrt{1 - \frac{1}{2Q^2}}$$

Les amplitudes sont respectivement pour ces deux pulsations :

$$X(0) = \frac{F}{k} = \frac{F}{m\omega_0^2}$$

et

$$X(\omega_m) = \frac{F}{m\omega_0^2} \frac{\omega_0^2}{2\xi\sqrt{\omega_0^2 - \xi^2}} = \frac{F}{m\omega_0^2} \frac{\omega_0^2\tau}{\sqrt{\omega_0^2 - \frac{1}{4\tau^2}}} = \frac{F}{m\omega_0^2} \frac{Q}{\sqrt{1 - \frac{1}{4Q^2}}}$$

On remarque que :

$$1 < \frac{Q}{\sqrt{1 - \frac{1}{4Q^2}}}$$

On aura donc un maximum pour $\omega = \omega_m$ et la valeur du maximum sera $\frac{F}{m\omega_0^2} \frac{\omega_0^2}{2\xi\sqrt{\omega_0^2 - \xi^2}}$. Il s'agit de la résonance en amplitude.

- si $\xi\sqrt{2} > \omega_0$ ou $Q = \omega_0\tau < \frac{1}{\sqrt{2}}$, alors il n'y a qu'une solution $\omega = 0$ et l'amplitude décroît avec la pulsation. Il n'y a pas de résonance.

On peut noter qu'il y a phénomène de résonance en amplitude pour un frottement peu important.

La dérivée de $X(\omega)$ étant nulle pour $\omega = 0$, la tangente à l'origine est horizontale.

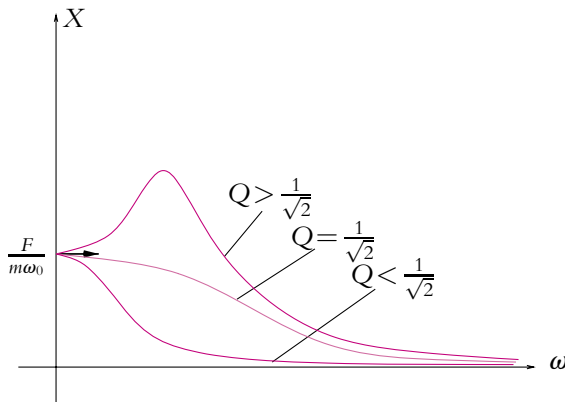


Figure 24.2 Résonance en amplitude.

Il s'agit de l'analogue de la résonance de la tension aux bornes du condensateur d'un circuit RLC série, étudiée en électrocinétique.

3.2 Bande passante

a) Calcul de la bande passante

La finesse de la résonance se définit à l'aide de la bande passante comme en électricité. Il s'agit du domaine de fréquences (ou de pulsations) pour lesquelles l'amplitude est au moins égale à l'amplitude maximale (amplitude à la résonance) divisée par $\sqrt{2}$. On a

$$X = \frac{F}{m\omega_0^2} \frac{1}{\sqrt{\left(1 - \left(\frac{\omega}{\omega_0}\right)^2\right)^2 + \frac{1}{Q^2} \frac{\omega^2}{\omega_0^2}}} \quad \text{et} \quad X_{\max} = \frac{F}{m\omega_0^2} \frac{Q}{\sqrt{1 - \frac{1}{4Q^2}}}$$

On cherche ω telles que

$$X \geq \frac{X_{\max}}{\sqrt{2}}$$

soit

$$\frac{F}{m\omega_0^2} \frac{1}{\sqrt{\left(1 - \left(\frac{\omega}{\omega_0}\right)^2\right)^2 + \frac{1}{Q^2} \frac{\omega^2}{\omega_0^2}}} \geq \frac{F}{m\omega_0^2} \frac{Q}{\sqrt{1 - \frac{1}{4Q^2}}} \frac{1}{\sqrt{2}}$$

La résolution de cette inégalité¹ conduit à :

$$\omega_1 = \omega_0 \sqrt{1 - \frac{1}{2Q^2} - \frac{1}{Q} \sqrt{1 - \frac{1}{4Q^2}}} \leq \omega \leq \omega_2 = \omega_0 \sqrt{1 - \frac{1}{2Q^2} + \frac{1}{Q} \sqrt{1 - \frac{1}{4Q^2}}}$$

soit une bande passante

$$\Delta\omega = \omega_2 - \omega_1$$

On note que ω_1 n'existe pas si $\frac{X_{\max}}{\sqrt{2}} > X(0)$ (c'est-à-dire si $Q > \frac{1}{2}\sqrt{4 + 2\sqrt{2}}$).

b) Cas d'un amortissement très faible

Dans ce cas, on a $Q \gg 1$ et on peut effectuer des développements limités des différentes quantités :

$$\omega_1 \simeq \omega_0 \left(1 - \frac{1}{2Q}\right) \quad \text{et} \quad \omega_2 \simeq \omega_0 \left(1 + \frac{1}{2Q}\right)$$

On en déduit la bande passante approchée :

$$\Delta\omega \simeq \frac{\omega_0}{Q}$$

¹ Pour le détail des calculs, on se reportera au cours d'électrocinétique, chapitre 15.

Ce résultat fournit une autre définition du facteur de qualité dans le cas où on a un faible amortissement :

$$Q = \frac{\omega_0}{\Delta\omega}$$

où ω_0 est la pulsation propre et $\Delta\omega$ la bande passante de la résonance.

La fréquence de résonance ω_m est pratiquement égale à ω_0 et l'amplitude maximale vaut en première approximation $X_{\max} = \frac{FQ}{m\omega_0}$. On en déduit pour un très faible amortissement que :

$$Q = \frac{X_{\max}}{X(0)}$$

On peut noter qu'il est assez difficile de faire les calculs de la bande passante dans le cas général, ce qui justifiera l'intérêt porté à la vitesse et à la puissance ou à l'énergie.

Cependant, physiquement, l'amplitude (ou l'élongation) est une grandeur importante qui ne doit pas dépasser une valeur critique sous peine de rupture du système. Par exemple, la rupture d'un pont peut être liée au passage d'une troupe marchant au pas si le pont présente la même fréquence propre que la cadence du pas. Il est d'ailleurs interdit aux troupes de marcher au pas sur un pont de manière à éviter la rupture du pont...

4. Résonance en vitesse

4.1 Expression de la vitesse

On a obtenu :

$$\underline{X} = \frac{\frac{F}{m}}{\omega_0^2 - \omega^2 + 2j\xi\omega}$$

On en déduit la vitesse :

$$\underline{V} = j\omega\underline{X} = \frac{j\omega\frac{F}{m}}{\omega_0^2 - \omega^2 + 2j\xi\omega} = \frac{\frac{FQ}{m\omega_0}}{1 + jQ\left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega}\right)}$$

En notation réelle, on a : $v(t) = V \cos(\omega t + \psi)$ avec :

$$\begin{cases} V = \omega X = \frac{\frac{F}{m}\omega}{\sqrt{(\omega_0^2 - \omega^2)^2 + 4\xi^2\omega^2}} = \frac{\frac{FQ}{m\omega_0}}{\sqrt{1 + Q^2\left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega}\right)^2}} \\ \psi = \varphi + \frac{\pi}{2} \end{cases}$$

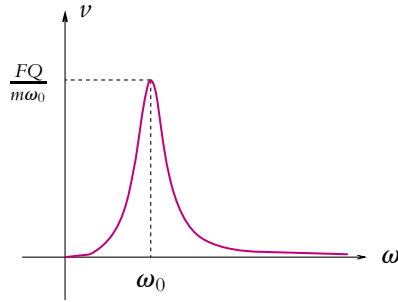


Figure 24.3 Résonance en vitesse.

4.2 Phénomène de résonance

La résonance en vitesse est obtenue quand l'amplitude de la vitesse V est maximale, c'est-à-dire quand $f(\omega) = 1 + Q^2 \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)^2$ est minimale, ce qui se produit quand $\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega}$ s'annule (puisque $f(\omega) \geq 1$). La résonance en vitesse a donc lieu pour $\omega = \omega_0$ et l'amplitude de la vitesse est alors :

$$V_{\max} = \frac{FQ}{m\omega_0}$$

4.3 Bande passante

On la définit toujours comme la gamme de fréquences pour laquelle on a une amplitude supérieure à l'amplitude maximale divisée par $\sqrt{2}$. Il faut donc résoudre

$$V \geq \frac{V_{\max}}{\sqrt{2}}$$

soit

$$\frac{\frac{FQ}{m\omega_0}}{\sqrt{1 + Q^2 \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)^2}} \geq \frac{FQ}{m\omega_0 \sqrt{2}}$$

ou encore :

$$Q^2 \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)^2 \leq 1$$

$$(\omega^2 - \omega_0^2)^2 \leq \frac{1}{Q^2} \omega_0^2 \omega^2$$

$$\left(\omega^2 - \frac{\omega_0}{Q} \omega - \omega_0^2 \right) \left(\omega^2 + \frac{\omega_0}{Q} \omega - \omega_0^2 \right) \leq 0$$



Il est important de noter la méthode de calcul : on a factorisé au lieu de développer l'expression obtenue.

Comme pour la résonance aux bornes de la résistance d'un circuit R, L, C série en électricité, l'allure de la courbe $V(\omega)$ montre que $\omega_1 \leq \omega \leq \omega_2$ où ω_1 et ω_2 sont solutions de l'équation $V = \frac{V_{\max}}{\sqrt{2}}$. Il suffit alors de résoudre $\omega^2 - \frac{\omega_0}{Q}\omega - \omega_0^2 = 0$ et $\omega^2 + \frac{\omega_0}{Q}\omega - \omega_0^2 = 0$. On obtient :

$$\omega_1 = -\frac{\omega_0}{2Q} \left(1 - \sqrt{1 + 4Q^2}\right) \leq \omega \leq \omega_2 = \frac{\omega_0}{2Q} \left(1 + \sqrt{1 + 4Q^2}\right)$$

On en déduit la bande passante :

$$\Delta\omega = \omega_2 - \omega_1 = \frac{\omega_0}{Q}$$

On en déduit une définition du facteur de qualité à partir de la résonance en vitesse :

$$Q = \frac{\omega_0}{\Delta\omega}$$

Il est à noter qu'ici la relation est stricte et non approchée comme dans le cas de l'étude de la position. Cette définition du facteur de qualité est donc toujours valable quelle que soit l'intensité de l'amortissement.

5. Impédance mécanique

On définit l'impédance mécanique d'un oscillateur comme le rapport de l'amplitude complexe de la force excitatrice par la vitesse complexe :

$$\underline{Z} = \frac{\underline{f}}{\underline{v}}$$

soit

$$\begin{aligned} \underline{Z} &= \frac{F}{j\omega \underline{X}} = \frac{F}{j\omega \frac{F}{k - m\omega^2 + j\omega\lambda}} \\ &= \lambda - j \left(\frac{k}{\omega} - m\omega \right) = m\omega_0 \left(\frac{1}{Q} - j \left(\frac{\omega_0}{\omega} - \frac{\omega}{\omega_0} \right) \right) \end{aligned}$$

On notera l'analogie avec l'électricité : l'intensité i correspond à la vitesse v , la tension excitatrice e à la force excitatrice f , l'inductance L à la masse m et l'inverse de la capacité $\frac{1}{C}$ à la raideur du ressort k et on a alors l'analogie entre impédance mécanique et impédance électrique.

6. Aspects énergétiques - Résonance en puissance

6.1 Expression de la puissance moyenne absorbée par l'oscillateur

La puissance instantanée absorbée par l'oscillateur est définie par :

$$\mathcal{P}(t) = \overrightarrow{f(t)} \cdot \overrightarrow{v(t)} = f(t)\dot{x}(t)$$

On obtient son expression en multipliant l'équation différentielle du mouvement $m\ddot{x} + \lambda\dot{x} + kx = f(t)$ par $\dot{x}$:

$$\mathcal{P}(t) = m\ddot{x}\dot{x} + \lambda\dot{x}^2 + kx\dot{x} = \frac{d}{dt} \left(\frac{1}{2}m\dot{x}^2 + \frac{1}{2}kx^2 \right) + \lambda\dot{x}^2$$

Le calcul de la valeur moyenne de $\mathcal{P}(t)$ sur une période T donne :

$$\begin{aligned} \mathcal{P} = \langle \mathcal{P}(t) \rangle &= \frac{1}{T} \int_{t_0}^{t_0+T} \frac{d}{dt} \left(\frac{1}{2}m\dot{x}^2 + \frac{1}{2}kx^2 \right) dt + \frac{1}{T} \int_{t_0}^{t_0+T} \lambda\dot{x}^2 dt \\ &= \frac{1}{T} \left[\frac{1}{2}m\dot{x}^2 + \frac{1}{2}kx^2 \right]_{t_0}^{t_0+T} + \lambda \langle \dot{x}^2 \rangle \end{aligned}$$

Comme x et ses dérivées sont de période T , $\left[\frac{1}{2}m\dot{x}^2 + \frac{1}{2}kx^2 \right]_{t_0}^{t_0+T} = 0$ et

$$\langle \mathcal{P}(t) \rangle = \lambda \langle \dot{x}^2 \rangle$$

6.2 Bilan énergétique

On peut établir ce résultat en appliquant directement le théorème de l'énergie cinétique :

$$d(E_c + E_p) = \delta W' = -\lambda\dot{x}^2 dt + f(t)\dot{x} dt$$

En moyenne sur une période, la variation d'énergie mécanique est nulle (pour le régime permanent qui est ici sinusoïdal forcé) donc :

$$\langle -\lambda\dot{x}^2 \rangle + \langle f(t)\dot{x} \rangle = 0$$

En moyenne, l'énergie apportée par la force $\overrightarrow{f}$ compense exactement l'énergie perdue par frottement.

6.3 Résonance en puissance

La recherche de la résonance en puissance consiste à déterminer les conditions permettant d'avoir une puissance moyenne absorbée par l'oscillateur maximale.

a) Phénomène de résonance

Comme $\dot{x}(t) = v(t) = V \cos(\omega t + \psi)$, on en déduit :

$$\begin{aligned}\langle \dot{x}^2 \rangle &= V^2 \langle \cos^2(\omega t + \psi) \rangle \\ &= \frac{V^2}{T} \int_{t_0}^{t_0+T} \cos^2(\omega t + \psi) dt \\ &= \frac{V^2}{T} \int_{t_0}^{t_0+T} \frac{1 + \cos(2\omega t + 2\psi)}{2} dt \\ &= \frac{V^2}{2}\end{aligned}$$

Par conséquent, en tenant compte de l'expression de V obtenue au paragraphe 4.1, on a :

$$\begin{aligned}\mathcal{P}(t) &= \lambda \frac{V^2}{2} = \frac{\xi F^2 \omega^2}{m \left((\omega_0^2 - \omega^2)^2 + 4\xi^2 \omega^2 \right)} \\ &= \frac{\omega_0 Q F^2 \omega^2}{2m \left(Q^2 (\omega_0^2 - \omega^2)^2 + \omega_0^2 \omega^2 \right)}\end{aligned}$$

ou en divisant par ω^2 :

$$\mathcal{P} = \frac{\xi F^2}{m \left(\left(\frac{\omega_0^2}{\omega} - \omega \right)^2 + 4\xi^2 \right)} = \frac{\omega_0 Q F^2}{2m \left(Q^2 \left(\frac{\omega_0^2}{\omega} - \omega \right)^2 + \omega_0^2 \right)}$$

Les variations de $\mathcal{P}$ en fonction de la pulsation ω sont inverses de celles de $f(\omega) = \left(\frac{\omega_0^2}{\omega} - \omega \right)^2 + 4\xi^2$. Or pour toute pulsation ω , $f(\omega) \leq 4\xi^2$ et $f(\omega_0) = 4\xi^2$. Donc f passe par un minimum pour $\omega = \omega_0$ et $\mathcal{P}$ admet un maximum pour $\omega = \omega_0$. La valeur du maximum est $\mathcal{P}_{\max} = \frac{F^2}{4\xi m} = \frac{Q F^2}{2m \omega_0^2}$.

On a donc résonance en puissance à la pulsation propre ω_0 .

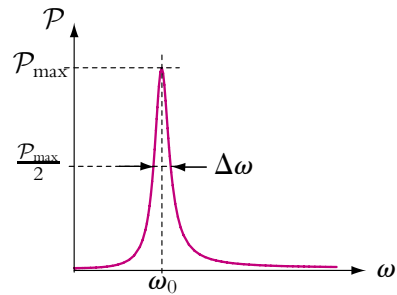


Figure 24.4 Résonance en puissance.

b) Bande passante de la résonance en puissance

Puisqu'il s'agit d'une quantité quadratique, la détermination de la bande passante s'obtient en cherchant les pulsations ω pour lesquelles on a :

$$\mathcal{P} \leq \frac{\mathcal{P}_{\max}}{2}$$

soit en résolvant :

$$\frac{\xi F^2}{m \left(\left(\frac{\omega_0^2}{\omega} - \omega \right)^2 + 4\xi^2 \right)} \leq \frac{F^2}{8\xi m}$$

ou après factorisation (analogue à celle faite pour la résonance en vitesse) :

$$(\omega_0^2 - 2\xi\omega - \omega^2)(\omega_0^2 + 2\xi\omega - \omega^2) \leq 0$$

En suivant la même démarche que pour la résonance en vitesse, on établit que les limites de la bande passante sont solutions du cas d'égalité soit :

$$\omega_1 = -\xi + \sqrt{\xi^2 + \omega_0^2} \quad \text{et} \quad \omega_2 = \xi + \sqrt{\xi^2 + \omega_0^2}$$

On en déduit que la bande passante vaut :

$$\Delta\omega = 2\xi = \frac{\omega_0}{Q}$$

C'est bien la même impression que pour la résonance en vitesse.

7. Portrait de phase d'un oscillateur forcé

Dans le cas des oscillations forcées, le second membre de l'équation différentielle est non nul. Ce terme impose la solution du régime permanent qui est un mouvement sinusoïdal à la pulsation d'excitation. Cette solution, analogue du point de vue de l'expression au cas de l'oscillateur harmonique, donne une trajectoire de phase elliptique pour le régime permanent.

Il y a également un second terme correspondant au régime transitoire qui s'exprime sous la forme d'une des solutions de l'oscillateur amorti et tend à disparaître au cours du temps.

On aura donc à distinguer deux parties dans la trajectoire de phase : la première du régime transitoire où coexistent les termes en importance sensiblement égale, la second, où ne reste plus que le terme du régime permanent.

Pour la première phase, on aura une convergence plus ou moins rapide suivant la nature de l'amortissement avec les mêmes distinctions que dans le cas de l'oscillateur amorti.

Par exemple, on pourra avoir des trajectoires de phase du type dans le cas d'un régime transitoire pseudo-périodique :

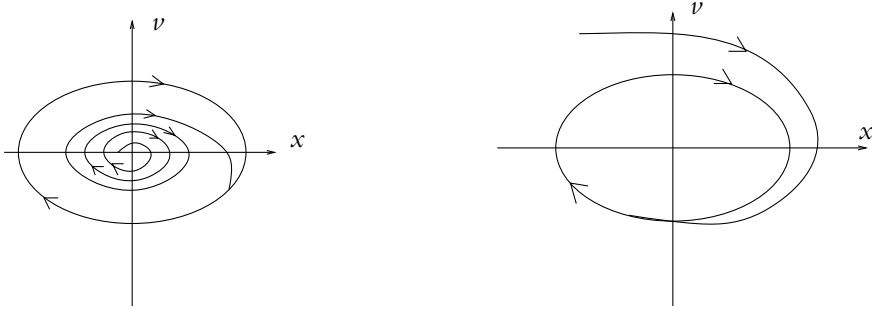


Figure 24.5 Portrait de phase d'un oscillateur forcé.

A. Applications directes du cours

1. Étude de l'élongation d'un oscillateur

L'oscillateur est un point M de masse m astreint à se déplacer le long d'un axe horizontal Ox . Le point M est attaché à une des extrémités d'un ressort (k, ℓ_0) , l'autre extrémité est reliée au point O . Un système d'amortissement exerce une force $\vec{F}_v = -h\dot{x}\vec{u}_x$. Le système est aussi soumis à une force excitatrice $\vec{F} = mF_0 \cos(\omega t)\vec{u}_x$, où F_0 est une constante positive.

1. Établir l'équation différentielle de la position x de M .

On effectue le changement de variable $X = x - x_e$ où x_e est la position d'équilibre de M . Montrer que l'équation différentielle peut alors se mettre sous la forme :

$$\ddot{X} + \frac{2}{\tau}\dot{X} + \omega_0^2 X = F_0 \cos(\omega t)$$

2. a) On cherche des solutions au régime forcé sous la forme $X = X_0 \cos(\omega t + \varphi)$. En utilisant la notation complexe déterminer $X_0(\omega)$ et $\varphi(\omega)$.

b) Représenter X_0 et φ en fonction de ω .

2. Vitesse en régime sinusoïdal forcé

On fait l'étude expérimentale d'un oscillateur grâce à un dispositif qui est sensible à la vitesse. C'est donc à la vitesse et non à l'élongation que l'on s'intéresse dans cet exercice.

L'oscillateur est un point M de masse m suspendu à un ressort (k, ℓ_0) dont la position est repérée par un axe Oz descendant. Un système d'amortissement exerce une force $\vec{F}_v = -h\dot{z}\vec{u}_z$. Le système est aussi soumis à une force excitatrice $\vec{F} = mF_0 \cos(\omega t)\vec{u}_z$, où F_0 est une constante positive. On pose $\omega_0^2 = k/m$ et $Q = m\omega_0/h$.

1. Établir l'équation différentielle à laquelle obéit la vitesse $v = \dot{z}$ en fonction de ω_0 , Q , F_0 et ω .

2. a) On cherche des solutions de la forme $v = V_0 \cos(\omega t + \varphi)$ pour le régime sinusoïdal forcé. On utilise alors des grandeurs complexes $\underline{V} = V_0 \exp(i(\omega t + \varphi))$. Déterminer l'expression de $\underline{V}$ en fonction des données puis celles de $V_0(\omega)$ et $\varphi(\omega)$.

b) Calculer la pulsation de résonance pour laquelle V_0 est maximale. Représenter $V_0(\omega)$ ainsi que $\varphi(\omega)$.

3. a) Exprimer la puissance cinétique $\mathcal{P}_c$ et la puissance de la force de tension du ressort $\mathcal{P}_r$ et montrer qu'elles sont nulles en moyenne.

b) Montrer qu'en moyenne la puissance de la force excitatrice $\mathcal{P}_e$ est égale à la puissance de la force de frottement $\mathcal{P}_f$. Montrer que cette puissance moyenne est maximale à la résonance de vitesse.

3. Pourquoi le ciel est-il bleu ? (Modèle de Thomson)

On considère un électron M d'une molécule dont le noyau est immobile en O . Cet électron est soumis à une force de rappel le liant à la molécule $\vec{F} = -k \vec{OM}$ ainsi qu'à une force de freinage $\vec{f}_v = -\lambda \vec{v}$ où $\vec{v}$ est la vitesse de M . Une onde lumineuse caractérisée par un champ électrique $\vec{E} = \vec{E}_0 \cos(\omega t)\vec{u}_x$ exerce la force $\vec{f}_e = -e\vec{E}$ sur l'électron. On néglige le poids.

- Déterminer l'équation différentielle du mouvement de l'électron.
- On s'intéresse au régime sinusoïdal forcé. Définir les grandeurs complexes appropriées à la résolution de l'équation différentielle. En déduire la solution forcée pour $\overrightarrow{OM}$.
- On s'intéresse à une molécule de l'atmosphère éclairée par la lumière du soleil : $m = 9,1 \cdot 10^{-31}$ kg, $e = 1,6 \cdot 10^{-19}$ C, $k = 100$ N.m⁻¹, $h = 10^{-20}$ kg.s⁻¹. Calculer numériquement w_0 (pulsation propre) et Q (facteur de qualité).
- La relation entre la longueur d'onde λ de l'onde lumineuse et ω est $\omega = \frac{2\pi c}{\lambda}$: avec $c = 3 \cdot 10^8$ m.s⁻¹.
Calculer les pulsations limites du visible et donner une expression approchée de $\overrightarrow{OM}$.
- Sachant que la puissance réémise par l'électron est proportionnelle au carré de son accélération, calculer le rapport émis pour le rouge ($\lambda = 800$ nm) et pour le bleu ($\lambda = 450$ nm). En déduire pourquoi le ciel apparaît bleu.

B. Exercices et problèmes

1. Résonance d'un système de deux points couplés par des ressorts

Deux points matériels A et B de même masse m sont reliés entre eux par un ressort de raideur K , et à deux points fixes par deux ressorts de raideur k . L'ensemble coulisse sans frottement sur une tige horizontale fixe. On note $\vec{u}$ un vecteur unitaire de cet axe.

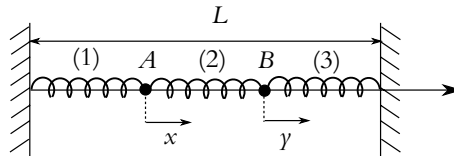


Figure 24.6

On note x et y les élongations de A et B comptées à partir de leur position d'équilibre.

- Écrire les relations à l'équilibre reliant les longueurs à vide des ressorts (ℓ_0 pour (1) et (3), L_0 pour (2)) et les longueurs à l'équilibre (ℓ_{eq} pour (1) et (3), L_{eq} pour (2)).
- Le point A subit une force supplémentaire $\vec{f} = mF_0 \cos(\omega t) \vec{u}$. Déterminer les équations du mouvement. Déduire des résultats précédents deux équations différentielles.
- a) On cherche pour x et y , des solutions au régime sinusoïdal forcé de la forme $x = X_0 \cos(\omega t + \varphi)$ et $y = Y_0 \cos(\omega t + \psi)$. On définit les grandeurs complexes associées $\underline{X} = X_0 e^{i(\omega t + \varphi)}$ et $\underline{Y} = Y_0 e^{i(\omega t + \psi)}$ avec $x = \text{Re}(\underline{X})$ et $y = \text{Re}(\underline{Y})$. Montrer que les solutions des équations différentielles précédentes sont :

$$\underline{X} = \frac{\omega_2^2 - \omega^2}{(\omega_0^2 - \omega^2)(\omega_1^2 - \omega^2)} F_0 e^{i\omega t} \quad \text{et} \quad \underline{Y} = \frac{\omega_3^2}{(\omega_0^2 - \omega^2)(\omega_1^2 - \omega^2)} F_0 e^{i\omega t}$$

où ω_0 , ω_1 , ω_2 et ω_3 sont des grandeurs à exprimer en fonction de k , K et m .

- b) Déduire des expressions précédentes $X_0(\omega)$, $Y_0(\omega)$, φ et ψ .
 c) Représenter $X_0(\omega)$, $Y_0(\omega)$ en fonction de ω . Pourquoi y-a-t-il des résonances infinies ?

2. Sismomètre pendule, d'après Banque PT 2005

Le sismomètre pendule est constitué d'une masse M reliée à un point A d'un châssis lui-même solidaire du sol en vibration par rapport à un référentiel $\mathcal{R}$ galiléen fixe. La liaison de M au bâti est modélisée par un comportement élastique de constante de raideur k associé à un frottement fluide caractérisé par la constante D selon le schéma de la figure (24.7).

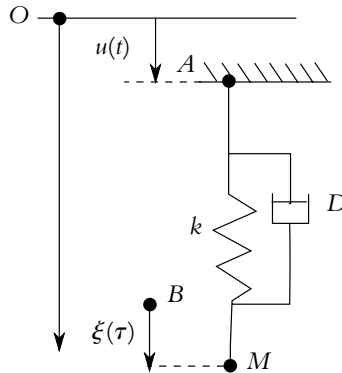


Figure 24.7

- Qu'appelle-t-on frottement fluide ou visqueux ?
- On note Oz un axe vertical descendant, O étant fixe dans $\mathcal{R}$. Le mouvement du sol par rapport au référentiel $\mathcal{R}$ est noté $u(t)$ (mouvement de A par rapport à O) et le mouvement de M par rapport au châssis (mouvement de M par rapport à B lié au châssis), donc par rapport au sol, est noté $\xi(t)$. On suppose que $\xi(t)$ est nul à l'équilibre en l'absence de tremblement de terre (donc lorsque u est constamment nul).
 - On pose : $\eta = \frac{D}{2m}$ et $\omega_s^2 = \frac{k}{M}$. En utilisant la relation de Chasles puis en la dérivant, montrer que :

$$\ddot{\xi} + 2\eta\dot{\xi} + \omega_s^2\xi = -\ddot{u}$$
 - Dans le cas de mouvements très rapides (donc de très grandes fréquences) quel est le terme prépondérant (au premier membre) ?
 Quelle grandeur représente alors $\xi(t)$?
 - Mêmes questions dans le cas de mouvements très lents.
- Déterminer l'expression complexe de la réponse $\xi(t)$, en régime permanent sinusoïdal (ou harmonique forcé) du système à un déplacement du sol de type sinusoïdal d'amplitude $u(t) = U_0 \cos(\omega t)$.
 - On caractérise le mouvement enregistré par son amplitude $X(\omega)$ et son déphasage $\Phi(\omega)$ et on introduit le paramètre $h = \frac{\eta}{\omega_s}$.

Déterminer les expressions de $X(\omega)$ et $\Phi(\omega)$ en fonction de h , U_0 et du rapport $x = \omega/\omega_s$.

c) Étudier les comportements limites de $X(\omega)$ et $\Phi(\omega)$ aux très basses, puis aux très hautes pulsations.

d) On pose $Y = \frac{U_0^2}{X^2}$. En exprimant Y en fonction de $x' = \frac{1}{x} = \frac{\omega_s}{\omega}$, montrer qu'il ne peut pas y avoir de résonance si h est supérieur à une valeur limite que l'on déterminera.

e) Tracer les courbes $X(\omega)$ et $\Phi(\omega)$ pour les deux valeurs suivantes du paramètre h : $h = 2$ puis $h = 0,5$, les valeurs de U_0 et de ω_s étant fixées.

Que deviennent ces courbes pour h nul ?

4. On définit la sensibilité σ d'un pendule sismique par la valeur absolue du quotient du déplacement $\Delta\xi$ de la masse M par la variation de l'accélération sismique Δa qui provoque ce déplacement.

a) Exprimer σ en fonction de la période propre du pendule dans le cas de mouvements dits de longue période (donc de mouvement très lents).

b) Application numérique : calculer le déplacement de la masse M , provoqué sur un pendule de période propre $T_0 = 60$ s par une variation d'accélération du bâti de $\Delta a = 10^{-3}$ cm.

3. Étude de la suspension d'un véhicule, d'après ENSTIM 2006

Le véhicule étudié est modélisé par un parallélépipède, de centre de gravité G et de masse M , reposant sur une roue par l'intermédiaire de la suspension dont l'axe OG reste toujours vertical. L'ensemble est animé d'une vitesse horizontale $\vec{v} = v\vec{u}_x$.

La suspension quant à elle, est modélisée par un ressort de constante de raideur $k = 1,0 \cdot 10^5$ N.m⁻¹ (de longueur à vide l_0) et un amortisseur fluide de constante d'amortissement constante $\lambda = 4,0 \cdot 10^3$ U.S.I. La masse de l'ensemble est $M = 1000$ kg. On note $\vec{g}$ l'accélération de la pesanteur uniforme et verticale.

La position verticale du véhicule est repérée par z_G dans le référentiel galiléen proposé ayant son origine sur la ligne moyenne des déformations du sol. On note z_O la cote du centre de la roue par rapport au niveau moyen de la route.

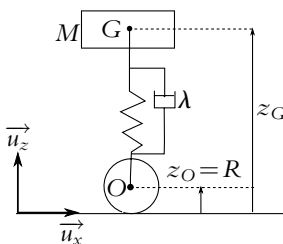


Figure 24.8

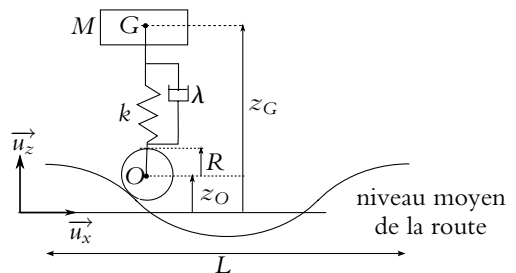


Figure 24.9

L'amortissement entre M et la roue introduit une force de frottement fluide, exercée par l'amortisseur sur M , qui s'écrit :

$$\vec{F} = -\lambda \left(\frac{dz_G}{dt} - \frac{dz_O}{dt} \right) \vec{u}_z$$

La route est parfaitement horizontale (Fig 24.8)

1. La route ne présente aucune ondulation et le véhicule n'a aucun mouvement vertical. Déterminer la position z_G de G lorsque le véhicule est au repos.

2. Suite à une impulsion soudaine, le véhicule acquiert un mouvement d'oscillations verticales. On cherche dans cette question à établir l'équation différentielle caractéristique du mouvement par une méthode énergétique.

On étudie le mouvement par rapport à la position d'équilibre établie précédemment. On posera $z = z_G - z_{G_{eq}}$.

a) Établir l'expression de l'énergie potentielle de pesanteur.

b) Établir l'expression de l'énergie potentielle élastique.

Les énergies potentielles seront exprimées en fonction de z et à une constante additive près.

c) Appliquer le théorème de l'énergie cinétique sous forme différentielle (ou théorème de la puissance cinétique) à la masse et en déduire l'équation différentielle du mouvement.

d) Établir l'équation différentielle du mouvement. Quel est le type de régime observé ?

e) Retrouver l'équation en appliquant directement la relation fondamentale de la dynamique.

f) Déterminer l'expression générale de la solution $z(t)$ en fonction de k , M et λ . On ne demande pas de déterminer les constantes d'intégration. Dessiner qualitativement l'allure de $z(t)$ sachant que $z(0) = 0$ et $\frac{dz}{dt}(0) > 0$.

La route est ondulée (Fig 24.9)

Le véhicule se déplace à vitesse horizontale constante v sur un sol ondulé. L'ondulation est assimilée à une sinusoïde de période spatiale L et d'amplitude A . La grandeur z_O peut alors s'écrire $z_O = R + A \cos \omega t$.

On étudie maintenant le mouvement par rapport à la position d'équilibre précédemment établie.

On posera $z = z_G - z_{G_{eq}}$.

Pour les applications numériques, on prendra $L = 1 \text{ m}$; $A = 10 \text{ cm}$.

3. Quelle est l'unité de λ ?

4. Exprimer ω en fonction de v et L .

5. En appliquant le principe fondamental de la dynamique à la masse M dans le référentiel terrestre supposé galiléen, établir l'équation différentielle en z régissant le mouvement : $m\ddot{z} + \lambda\dot{z} + kz = k(z_0 - R) + \lambda\dot{z}_0$.

6. Justifier qualitativement le fait que l'on recherche la solution $z(t)$ de cette équation différentielle sous une forme sinusoïdale $z(t) = z_{\max} \cos(\omega t + \varphi)$.

7. Résolution par la méthode des complexes

On pose $z_O - R = \underline{A} \exp(j\omega t)$ et $\underline{z} = \underline{Z} \exp(j\omega t)$, réponse complexe du véhicule à l'excitation sinusoïdale.

- a) Montrer que $\frac{Z}{A} = \frac{\left(\frac{k}{M} + j\frac{\omega\lambda}{M}\right)}{\left(-\omega^2 + j\omega\frac{\lambda}{M} + \frac{k}{M}\right)}$ avec j tel que $j^2 = -1$ puis que l'on peut mettre sous la forme

$$\frac{Z}{A} = \frac{1 + j\frac{\omega}{\omega_1}}{1 - \frac{\omega^2}{\omega_0^2} + j\frac{\omega}{Q\omega_0}} = \underline{H_1} \times \underline{H_2}$$

Exprimer alors ω_0 , ω_1 et Q en fonction de k , λ et M .

- b) Calculer numériquement ω_0 , ω_1 et Q .

- c) Donner l'expression du module $\left|\frac{Z}{A}\right|$ en fonction de ω_0 , ω_1 et Q .

8. Étude fréquentielle

On souhaite étudier l'amplitude des oscillations en fonction de la vitesse de la voiture. Pour cela on étudie $\left|\frac{Z}{A}\right|$ sous la forme $|\underline{H_1} \underline{H_2}|$ en fonction de ω .

- a) Tracer le diagramme de Bode asymptotique relatif à $\left|\frac{Z}{A}\right|$. Tracer l'allure de $\left|\frac{Z}{A}\right|$.

Remarque : on pourra d'abord tracer les diagrammes de $|\underline{H_1}|$ et $|\underline{H_2}|$ en fonction de ω .

- b) ω_r , la valeur de ω pour laquelle l'amplitude est maximale est de l'ordre de ω_0 . Quelle est la valeur de ν correspondante ? Calculer l'amplitude des oscillations du véhicule pour $\omega = \omega_0$.

9. Application

Dans le film "le salaire de la peur", Yves Montand conduit un camion ($\omega_0 \simeq 25 \text{ rad.s}^{-1}$) chargé de nitroglycérine. Il passe sur une "tôle ondulée" de période spatiale 1 m et pour laquelle $A = 10 \text{ cm}$. Afin d'éviter l'explosion du chargement, il doit traverser la tôle à une vitesse inférieure à 5 km.h^{-1} ou supérieure à 50 km.h^{-1} . Justifier qualitativement ceci à l'aide des résultats précédents.

4. Production de vagues dans une piscine, d'après ENSTIM 2003

Pour créer des vagues dans une piscine, on fait effectuer des oscillations verticales à une masse m immergée sur un côté du bassin.

Soit une masse M homogène de masse volumique ρ et de volume V plongée dans l'eau de masse volumique ρ_{eau} . Cette masse est suspendue à un ressort de raideur k et de longueur à vide l_0 accroché à son autre extrémité en un point A . On note Oz l'axe vertical passant par A orienté vers le bas et on prend A comme origine. A l'équilibre, la masse M est située en $z = h$. On se place dans le référentiel terrestre supposé galiléen.

- Établir la condition d'équilibre de la masse M dans l'air.
- En déduire l'équation différentielle du mouvement en z de la masse M si le mouvement s'effectue dans l'air.
- Quelle est la nature du mouvement ? On donnera les expressions de ses principales caractéristiques.

4. Du fait que le mouvement a lieu dans l'eau, comment doit-on modifier les équations précédentes ?
5. On tient compte désormais d'une force de frottement visqueux $\vec{f} = -\alpha \vec{v}$. Déterminer la nouvelle équation différentielle vérifiée par z .
6. Dans le cas d'un amortissement faible et en supposant que $z(0) = h_1 > h$ et $\dot{z}(0) = 0$, quelle est l'allure de $z(t)$?
7. A l'aide d'un piston, on impose un mouvement sinusoïdal au point de suspension A du ressort. Cela revient à appliquer une force $M\omega^2 a \cos \omega t \vec{u}_z$ à la masse M . Donner la nouvelle équation différentielle du mouvement.
8. Déterminer l'expression de l'amplitude Z des oscillations en fonction de a , $x = \frac{\omega}{\omega_0}$ et $\tau = \frac{M}{\alpha}$ avec $\omega_0 = \sqrt{\frac{k}{M}}$.
9. Déterminer les deux conditions nécessaires pour qu'on puisse avoir des oscillations d'amplitude supérieure à a , l'une portant sur la valeur minimale de la masse M et l'autre sur l'intervalle de pulsations à utiliser.
10. Quelle est alors la pulsation pour laquelle l'amplitude est la plus grande ?
11. Donner l'expression de l'amplitude correspondante en fonction de a , M , k et α .

Théorème du moment cinétique

25

Jusqu'à présent, on dispose de deux méthodes pour obtenir l'équation du mouvement d'un point matériel : l'utilisation du principe fondamental de la dynamique et celle du théorème de l'énergie cinétique, cette deuxième approche n'étant efficace que pour un système à un seul paramètre de position et avec des forces dont on sait déterminer le travail. L'objet de ce chapitre est d'introduire la notion de moment cinétique et d'établir une nouvelle méthode de résolution.

1. Définition du moment cinétique d'un point matériel

1.1 Moment cinétique par rapport à un point

Le *moment cinétique* par rapport à un point est une grandeur dynamique utile dans de nombreux cas par la simplicité qu'elle apporte au niveau de la résolution des problèmes.

Soit un point O du référentiel $\mathcal{R}$ dans lequel on travaille.

Soit un point matériel M de masse m , de vitesse $\vec{v}_{/\mathcal{R}}(M)$ et de quantité de mouvement $\vec{p}_{/\mathcal{R}}(M)$ par rapport au référentiel $\mathcal{R}$.

Le moment cinétique de M par rapport à O dans le référentiel $\mathcal{R}$ est défini par le vecteur :

$$\vec{L}_{O/\mathcal{R}}(M) = \vec{OM} \wedge \vec{p}_{/\mathcal{R}}(M) = \vec{OM} \wedge m \vec{v}_{/\mathcal{R}}(M)$$

► Remarque :

Le fait que le mouvement d'un point matériel soit déterminé par la donnée d'un seul vecteur implique qu'en mécanique du point, il n'est pas nécessaire d'utiliser en même temps quantité de mouvement et moment cinétique. Le problème posé sera résolu par la détermination de l'une ou de l'autre de ces quantités.

Là réside une différence essentielle avec la mécanique des solides où la détermination des deux quantités que sont le moment cinétique et la quantité de mouvement sera nécessaire.

1.2 Dépendance du moment cinétique par rapport au point considéré

Le moment cinétique dépend du point par rapport auquel on le détermine.

$$\begin{aligned}\vec{L}_{O'/\mathcal{R}}(M) &= \vec{O'M} \wedge \vec{p}_{/\mathcal{R}}(M) \\ &= (\vec{O'O} + \vec{OM}) \wedge \vec{p}_{/\mathcal{R}}(M) \\ &= \vec{O'O} \wedge \vec{p}_{/\mathcal{R}}(M) + \vec{OM} \wedge \vec{p}_{/\mathcal{R}}(M) \\ &= \vec{O'O} \wedge \vec{p}_{/\mathcal{R}}(M) + \vec{L}_{O/\mathcal{R}}\end{aligned}$$

En simplifiant les notations on a donc la relation :

$$\vec{L}_{O'/\mathcal{R}} = \vec{L}_{O/\mathcal{R}} + \vec{O'O} \wedge \vec{p}_{/\mathcal{R}}$$

On ne retiendra pas cette relation mais on la ré-établira au besoin et on se souviendra du fait que le moment cinétique dépend du point où on le détermine.

1.3 Moment cinétique par rapport à un axe

Soit un axe Δ . On note O un point de cet axe et $\vec{u}_\Delta$ un vecteur unitaire directeur de l'axe Δ .

On définit un moment cinétique L_Δ par rapport à l'axe Δ passant par O à l'aide de la relation :

$$L_\Delta = \vec{L}_{O/\mathcal{R}} \cdot \vec{u}_\Delta$$

Il est important de noter qu'un moment par rapport à un point est un vecteur tandis qu'un moment par rapport à un axe est un scalaire.

Cette notion n'aura un sens que si la quantité ainsi définie est indépendante du point O de l'axe. On peut le montrer en utilisant un autre point O' qui permet de définir un autre moment par rapport à l'axe :

$$L'_\Delta = \vec{L}_{O'/\mathcal{R}} \cdot \vec{u}_\Delta$$

On a établi que :

$$\vec{L}_{O'/\mathcal{R}} = \vec{L}_{O/\mathcal{R}} + \vec{O'O} \wedge \vec{p}_{/\mathcal{R}}$$

Alors :

$$L'_\Delta = (\vec{L}_{O/\mathcal{R}} + \vec{O'O} \wedge \vec{p}_{/\mathcal{R}}) \cdot \vec{u}_\Delta = \vec{L}_{O/\mathcal{R}} \cdot \vec{u}_\Delta + (\vec{O'O} \wedge \vec{p}_{/\mathcal{R}}) \cdot \vec{u}_\Delta = L_\Delta$$

car $\vec{O'O}$ et $\vec{u}_\Delta$ sont colinéaires, $\vec{OO'} \wedge \vec{p}_{/\mathcal{R}}$ est perpendiculaire à $\vec{O'O}$ et le produit scalaire de deux vecteurs perpendiculaire est nul. On parlera donc de *moment cinétique par rapport à l'axe Δ* .

2. Moment d'une force

2.1 Moment d'une force par rapport à un point

On appelle *moment de la force $\vec{f}$ par rapport à un point O* le vecteur :

$$\overrightarrow{\mathcal{M}}_O(\vec{f}) = \overrightarrow{OM} \wedge \vec{f}$$

où M désigne le point matériel sur lequel s'applique la force $\vec{f}$.

► **Remarque :** Par un raisonnement analogue à celui effectué pour le moment cinétique, on peut établir que le moment des forces dépend du point où on le calcule et que :

$$\overrightarrow{\mathcal{M}}_{O'}(\vec{f}) = \overrightarrow{\mathcal{M}}_O(\vec{f}) + \overrightarrow{O'O} \wedge \vec{f}$$



Remarque sur le calcul effectif du moment d'une force

Il s'agit d'une remarque d'ordre géométrique. Le moment par rapport à un point O d'une force $\vec{f}$ appliquée en un point A est par définition perpendiculaire à $\overrightarrow{OA}$ et à $\vec{f}$ du fait du produit vectoriel. Son module est égal au produit du module de la force par la distance du point O à la droite d'action de la force (cette distance est appelée *bras de levier*) :

$$\|\overrightarrow{OA} \wedge \vec{f}\| = \|\overrightarrow{OA}\| \times \|\vec{f}\| \sin \theta$$

où θ désigne l'angle entre OA et la droite d'action de $\vec{f}$. Dans le triangle rectangle OAM , on a :

$$|\sin \theta| = \frac{\|\overrightarrow{OM}\|}{\|\overrightarrow{OA}\|} = \frac{d}{\|\overrightarrow{OA}\|}$$

et

$$\|\overrightarrow{OA}\| \sin \theta = d$$

soit

$$\|\overrightarrow{OA} \wedge \vec{f}\| = d \|\vec{f}\|$$

ce qui fournit une méthode de calcul du moment d'une force.

Pour le sens du moment, on utilise le fait que $(\overrightarrow{OA}, \vec{f}, \overrightarrow{\mathcal{M}}_O(\vec{f}))$ forment une base directe.

En pratique, on se demande si $\vec{f}$ a tendance à faire tourner le point A dans le sens positif ou négatif. Dans le premier cas, $\overrightarrow{\mathcal{M}}_O$ est dirigé dans le sens de $\vec{u}_z$ et dans le second dans le sens de $-\vec{u}_z$.

Ce calcul revient à considérer soit le produit du module de la force par la composante perpendiculaire à la force de $\overrightarrow{OM}$, soit le produit de la distance OM par la composante perpendiculaire à $\overrightarrow{OM}$ de la force.

Il ne faut cependant pas confondre cette méthode de calcul du moment d'une force avec le moment d'une force par rapport à un axe qui sera détaillé dans le paragraphe suivant.

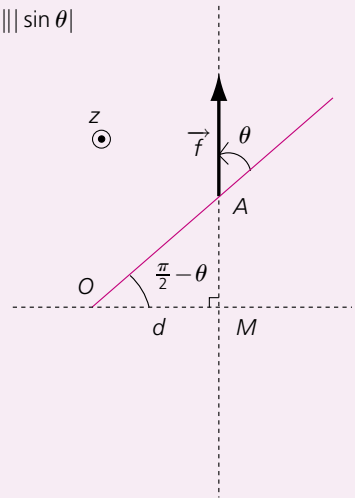


Figure 25.1 Calcul du moment d'une force.

2.2 Moment d'une force par rapport à un axe

Soit un axe Δ . On note O un point de cet axe et $\vec{u}_\Delta$ un vecteur directeur unitaire de cet axe.

Le moment résultant des forces s'exerçant sur un point matériel par rapport à l'axe Δ est par définition le scalaire :

$$\mathcal{M}_\Delta = \overrightarrow{\mathcal{M}}_O(\vec{f}) \cdot \vec{u}_\Delta$$

où $\overrightarrow{\mathcal{M}}_O(\vec{f})$ est le moment résultant des forces par rapport à O un point de l'axe Δ .

$\mathcal{M}_\Delta$ est indépendant du choix du point O de l'axe Δ (cette propriété s'établit de la même manière que pour le moment cinétique par rapport à un axe) : on parle donc de *moment de la force par rapport à l'axe Δ* .

3. Théorème du moment cinétique en référentiel galiléen

L'objet de ce paragraphe est d'établir le théorème du moment cinétique à partir du principe fondamental de la dynamique. Il ne s'agit donc que de formuler différemment ce principe afin de le rendre plus facile d'emploi dans certaines situations.

Toutes les quantités considérées sont définies par rapport à un référentiel galiléen.

3.1 Théorème du moment cinétique par rapport à un point fixe O

L'énoncé du théorème est le suivant :

La dérivée du moment cinétique du point M en un point fixe O par rapport au temps est égale à la somme des moments des forces appliquées en M par rapport au point fixe O .

Ceci se traduit mathématiquement par :

$$\frac{d\overrightarrow{L}_{O/\mathcal{R}}(M)}{dt} = \sum_i \overrightarrow{\mathcal{M}}_O(\vec{f}_i)$$

Ce théorème est une conséquence directe du principe fondamental de la dynamique.

En effet, en appliquant la définition du moment cinétique, on obtient :

$$\frac{d\overrightarrow{L}_{O/\mathcal{R}}(M)}{dt} = \frac{d}{dt} \left(\overrightarrow{OM} \wedge m\vec{v}_{/\mathcal{R}}(M) \right)$$

soit en dérivant le produit vectoriel :

$$\frac{d\vec{L}_{O/\mathcal{R}}(M)}{dt} = \frac{d\vec{OM}}{dt} \wedge m\vec{v}_{/\mathcal{R}}(M) + \vec{OM} \wedge \frac{d(m\vec{v}_{/\mathcal{R}}(M))}{dt}$$

Or, O étant fixe, $\frac{d\vec{OM}}{dt} = m\vec{v}_{/\mathcal{R}}(M)$.

En utilisant le principe fondamental de la dynamique, on déduit :

$$\frac{d\vec{L}_{O/\mathcal{R}}(M)}{dt} = \vec{v}_{/\mathcal{R}}(M) \wedge m\vec{v}_{/\mathcal{R}}(M) + \vec{OM} \wedge \sum_i \vec{f}_i$$

Or $\vec{v}_{/\mathcal{R}}(M) \wedge \vec{v}_{/\mathcal{R}}(M) = \vec{0}$ et $\vec{OM} \wedge \sum_i \vec{f}_i = \sum_i \vec{OM} \wedge \vec{f}_i$ donc :

$$\frac{d\vec{L}_{O/\mathcal{R}}(M)}{dt} = \sum_i \vec{OM} \wedge \vec{f}_i$$

Ce qui donne, avec la définition du moment d'une force par rapport à un point :

$$\frac{d\vec{L}_{O/\mathcal{R}}(M)}{dt} = \sum_i \vec{\mathcal{M}}_O(\vec{f}_i)$$

3.2 Cas où le point O n'est pas fixe

Dans ce cas, il est toujours possible d'introduire un point A fixe dans le référentiel galiléen considéré.

On décompose : $\vec{OM} = \vec{OA} + \vec{AM}$. En considérant toujours toutes les quantités définies par rapport au référentiel galiléen, la démonstration précédente devient alors :

$$\begin{aligned} \frac{d\vec{L}_{O/\mathcal{R}}(M)}{dt} &= \frac{d\vec{OM}}{dt} \wedge m\vec{v}_{/\mathcal{R}}(M) + \vec{OM} \wedge \frac{d(m\vec{v}_{/\mathcal{R}}(M))}{dt} \\ &= \frac{d\vec{AM}}{dt} \wedge m\vec{v}_{/\mathcal{R}}(M) - \frac{d\vec{AO}}{dt} \wedge m\vec{v}_{/\mathcal{R}}(M) + \vec{OM} \wedge \frac{d(m\vec{v}_{/\mathcal{R}}(M))}{dt} \\ &= \vec{v}_{/\mathcal{R}}(M) \wedge m\vec{v}_{/\mathcal{R}}(M) - \vec{v}_{/\mathcal{R}}(O) \wedge m\vec{v}_{/\mathcal{R}}(M) + \vec{OM} \wedge \sum_i \vec{f}_i \\ &= \vec{0} - \vec{v}_{/\mathcal{R}}(O) \wedge m\vec{v}_{/\mathcal{R}}(M) + \sum_i \vec{OM} \wedge \vec{f}_i \\ &= -\vec{v}_{/\mathcal{R}}(O) \wedge m\vec{v}_{/\mathcal{R}}(M) + \sum_i \vec{\mathcal{M}}_O(\vec{f}_i) \\ &= \sum_i \vec{\mathcal{M}}_O(\vec{f}_i) + m\vec{v}_{/\mathcal{R}}(M) \wedge \vec{v}_{/\mathcal{R}}(O) \end{aligned}$$

Lorsque le point O n'est pas fixe, il faut tenir compte d'un terme supplémentaire :

$$m\vec{v}_{/\mathcal{R}}(M) \wedge \vec{v}_{/\mathcal{R}}(O)$$

qui ne s'annule généralement pas. Cela complique habituellement la résolution et on évitera donc de choisir un point O qui n'est pas fixe. Les seuls cas où ce terme disparaît sont :

- si les vitesses des points O et M sont colinéaires, auquel cas le produit vectoriel est nul,
- si la vitesse du point O est nul, à savoir le cas où O est fixe,
- si le point M est immobile, ce qui *a priori* a peu de chances de se produire sauf si on se trouve à l'équilibre.



On n'appliquera le théorème du moment cinétique qu'en un point fixe O .

Ce théorème est particulièrement efficace lorsque la somme des moments des forces est nulle, ce qui est, par exemple, le cas lorsque toutes les forces ont une droite d'action passant par un même point fixe et qu'on choisit ce point comme référence pour le moment cinétique. On verra lors de l'étude du mouvement à force centrale (qui correspond au cas qui vient d'être décrit) le grand intérêt de ce théorème.

3.3 Théorème du moment cinétique par rapport à un axe

On rappelle que le moment par rapport à un axe Δ est le produit scalaire entre le moment par rapport à un point de cet axe et un vecteur unitaire directeur de cet axe. De ce fait, le théorème du moment cinétique par rapport à un axe se déduit du même théorème par rapport à un point. Dans un référentiel galiléen :

$$\frac{d\vec{L}_{O/\mathcal{R}}}{dt} = \sum_i \vec{\mathcal{M}}_O(\vec{f}_i)$$

qu'on multiplie scalairement par $\vec{u}_\Delta$ un vecteur directeur unitaire de Δ l'axe fixe considéré :

$$\frac{d\vec{L}_{O/\mathcal{R}}}{dt} \cdot \vec{u}_\Delta = \left(\sum_i \vec{\mathcal{M}}_O(\vec{f}_i) \right) \cdot \vec{u}_\Delta = \sum_i \mathcal{M}_\Delta(\vec{f}_i)$$

Comme $\vec{u}_\Delta$ est un vecteur constant puisque l'axe est fixe, on a :

$$\frac{d\vec{L}_{O/\mathcal{R}}}{dt} \cdot \vec{u}_\Delta = \frac{d\vec{L}_{O/\mathcal{R}} \cdot \vec{u}_\Delta}{dt} = \frac{dL_\Delta}{dt}$$

On a ainsi le théorème du moment cinétique par rapport à un axe fixe Δ en référentiel galiléen :

$$\frac{dL_{\Delta}}{dt} = \sum_i \mathcal{M}_{\Delta}(\vec{f}_i)$$

4. Exemple d'utilisation du théorème du moment cinétique : le pendule simple

On s'intéresse au mouvement d'un pendule simple c'est-à-dire au mouvement d'une masse ponctuelle m suspendue à l'extrémité d'un fil inextensible, de masse négligeable, de longueur l dont l'autre extrémité est fixe.

À l'instant initial, le pendule est lâché sans vitesse d'une position définie par l'angle θ_0 entre la verticale descendante et la direction du pendule.

On suit l'approche en quatre points pour résoudre un problème de mécanique :

1. Système : masse m .
2. Référentiel du laboratoire considéré comme galiléen.
3. Bilan des forces :
 - poids $m\vec{g}$,
 - tension exercée par le fil $\vec{R}$.
4. Résolution :

On va appliquer la méthode présentée dans ce chapitre en utilisant le théorème du moment cinétique et on donnera également la résolution par le principe fondamental de la dynamique et par une approche énergétique pour montrer que les trois approches sont possibles.

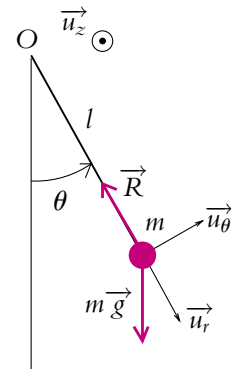


Figure 25.2 Pendule simple.

4.1 Résolution par le théorème du moment cinétique

Avant d'appliquer le théorème du moment cinétique, on remarque que le mouvement est plan. En effet, la projection du principe fondamental de la dynamique sur la direction orthogonale au plan défini par l'axe vertical (Oz) et la position initiale du pendule donne : $\ddot{z} = 0$ donc $\dot{z} = 0$ et $z = 0$ compte tenu des conditions initiales : le mouvement du pendule est plan.

Le théorème du moment cinétique s'écrit :

$$\frac{d\vec{L}_O}{dt} = \vec{\mathcal{M}}_O(m\vec{g}) + \vec{\mathcal{M}}_O(\vec{R})$$

On calcule en $\vec{L}_O$ coordonnées cylindriques :

$$\vec{L}_O = \vec{OM} \wedge m \vec{v} = \begin{vmatrix} l & 0 \\ 0 & ml\dot{\theta} \\ 0 & 0 \end{vmatrix} \wedge \begin{vmatrix} 0 \\ ml\dot{\theta} \\ 0 \end{vmatrix} = \begin{vmatrix} 0 \\ 0 \\ ml^2\dot{\theta} \end{vmatrix}$$

Donc :

$$\frac{d\vec{L}_O}{dt} = \frac{d}{dt} (ml^2\dot{\theta}\vec{u}_z) = ml^2\ddot{\theta}\vec{u}_z$$

La droite d'action de $\vec{R}$ passe par O donc $\vec{\mathcal{M}}_O(\vec{R}) = \vec{0}$.

Quant au poids, son moment par rapport à O est égal au produit de mg par la distance de la droite d'action de $m\vec{g}$ à l'axe Oz c'est-à-dire $l \sin \theta$. Pour connaître le signe du moment, on regarde dans quel sens le poids a tendance à faire tourner le point. Ici, c'est dans le sens négatif de $\vec{u}_z$ donc :

$$\vec{\mathcal{M}}_0(m\vec{g}) = -mgl \sin \theta \vec{u}_z$$

On note qu'il est également possible d'obtenir ce résultat par un calcul direct :

$$\vec{\mathcal{M}}_0(m\vec{g}) = \vec{OM} \wedge m\vec{g} = \begin{vmatrix} l & mg \cos \theta \\ 0 & -mg \sin \theta \\ 0 & 0 \end{vmatrix} \wedge \begin{vmatrix} 0 \\ 0 \\ -mgl \sin \theta \end{vmatrix}$$

On obtient finalement :

$$ml^2\ddot{\theta} = -mgl \sin \theta$$

et l'équation du mouvement est :

$$\ddot{\theta} + \frac{g}{l} \sin \theta = 0$$

La résolution de cette équation différentielle non linéaire (du fait de la présence du terme $\sin \theta$) n'est pas simple et on ne détaillera pas plus ici le mouvement.

4.2 Résolution par le principe fondamental de la dynamique

Le principe fondamental de la dynamique s'écrit :

$$m\vec{a} = m\vec{g} + \vec{R}$$

On se place en coordonnées polaires dans le plan vertical où s'effectue le mouvement du pendule. L'expression de l'accélération est :

$$\vec{a} = -l\dot{\theta}^2\vec{u}_r + l\ddot{\theta}\vec{u}_\theta$$

en tenant compte du fait que $r = l$ est constant.

On projette le principe fondamental de la dynamique suivant $\vec{u}_\theta$: on s'affranchit ainsi de la tension $\vec{R}$ qui n'a de composante non nulle que sur $\vec{u}_r$. On obtient :

$$ml\ddot{\theta} = -mg \sin \theta$$

soit

$$\ddot{\theta} + \frac{g}{l} \sin \theta = 0$$

ce qui est bien la même équation que celle que l'on obtient en appliquant le théorème du moment cinétique en O .

4.3 Méthode énergétique

La réaction $\vec{R}$ est perpendiculaire au mouvement et ne travaille donc pas. Le poids est une force conservative. Par conséquent, on a conservation de l'énergie mécanique soit, en explicitant son expression dans le système de coordonnées choisi :

$$Em = Ep + Ec = -mgl \cos \theta + C + \frac{1}{2} m (\dot{l}\dot{\theta})^2$$

en notant C la constante de l'énergie potentielle de pesanteur. On note qu'il n'est pas nécessaire de préciser la valeur de C puisqu'on va dériver l'expression de l'énergie mécanique pour obtenir l'équation du mouvement et que la dérivée d'une constante est nulle. Par dérivation, on obtient :

$$ml\dot{\theta}\ddot{\theta} + mgl \sin \theta \dot{\theta} = 0$$

Comme on considère le cas du mouvement, le cas $\dot{\theta} = 0$ est inintéressant. Par conséquent, on en déduit après simplification :

$$\ddot{\theta} + \frac{g}{l} \sin \theta = 0$$

ce qui est encore une fois la même équation.

4.4 Quelle méthode choisir ?

La bonne méthode en mécanique du point est celle qui élimine les forces inconnues (les tensions de fils et les réactions de support essentiellement). Dans le cas présent, la projection du principe fondamental de la dynamique sur une direction orthogonale à $\vec{R}$ ou le théorème du moment cinétique ou une approche énergétique remplissent ce critère : les trois méthodes sont bonnes, aucune n'apparaît préférable à l'autre.

A. Applications directes du cours

1. Véhicule sur une colline

Un véhicule M de masse m se déplace de haut en bas d'une colline ; la trajectoire est assimilée à un quart de cercle vertical de centre O et de rayon R . On note θ , l'angle que fait OM avec la verticale. Il démarre du point le plus haut ($\theta = 0$) avec une vitesse v_0 . Il n'y a pas de frottement.

1. En appliquant le théorème du moment cinétique par rapport à O , déterminer une équation différentielle en fonction de θ et ses dérivées.

2. Intégrer cette équation après l'avoir multipliée par $\frac{d\theta}{dt}$, et déterminer l'expression de $\frac{d\theta}{dt}$.

2. Modèle de l'atome d'hydrogène

Soit un proton O de charge $+e$ fixe, et un électron M de charge $-e$ en mouvement circulaire uniforme (rayon R) autour du proton. L'expérience montre que l'énergie totale du système est quantifiée, c'est-à-dire qu'elle peut s'écrire sous la forme :

$$E_m = -\frac{k}{n^2}$$

où n est un entier naturel et k une constante positive ne dépendant que des caractéristiques du système.

1. Calculer la force s'exerçant sur M et l'énergie de M en coordonnées polaires de centre O . En déduire la vitesse v et le rayon R en fonction de ϵ_0 , k , e et n .

2. Calculer le moment cinétique $\vec{L}_O$ de l'électron par rapport à O . Montrer qu'il est lui-même quantifié de la forme $L_O = nC$ où C est une constante à déterminer

3. Point matériel lié à un fil

Un point matériel M de masse m se déplace sans frottement sur un plan horizontal. Il est attaché à un fil de masse négligeable. Le plan est percé au point O choisi pour origine, le fil traverse le plan en O ; un opérateur exerce sur l'autre extrémité du fil une force de traction de module F constant.

À $t = 0$, le point M se trouve sur l'axe Ox à une distance D de O , sa vitesse est alors V_0 , colinéaire à Oy .

1. Montrer que le moment cinétique $\vec{L}_O$ en O de M se conserve pendant toute la durée de la traction. Calculer son module et sa direction. Définir la vitesse aréolaire et montrer que le mouvement de M se fait suivant la loi des aires (voir chapitre 30, paragraphe 4.5).

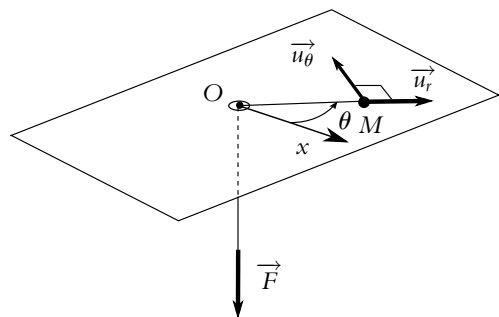


Figure 25.3

2. Déterminer l'énergie mécanique E_m de M ; on montrera que la force de traction dérive d'une énergie potentielle E_p que l'on exprimera en utilisant les coordonnées cylindriques. On prendra $E_p = 0$ pour $OM = 0$.
3. a) Montrer que l'on peut se ramener à un problème à une dimension et définir la fonction énergie potentielle efficace.
b) Étudier ses variations en fonction de $r = OM$. Montrer qu'elle passe par un minimum pour $r_1 = \left(\frac{mV_0^2 D^2}{F} \right)^{1/3}$.
c) Tracer le graphe correspondant et faire apparaître l'énergie mécanique. Déterminer graphiquement les deux distances extrêmes de M à O ; à quelle condition sur F la position initiale est-elle le point de la trajectoire le plus éloigné de O ? Le point matériel peut-il arriver jusqu'à O ?
d) Quelle valeur faut-il donner à F pour observer un mouvement circulaire ? Montrer qu'alors il est nécessairement uniforme.

4. Perle sur un cerceau

Une perle de masse m peut se déplacer sans frottement sur un cerceau de rayon r . On suppose que le référentiel terrestre est galiléen. Établir l'équation différentielle du mouvement de la perle en appliquant le théorème du moment cinétique.

5. Saut à ski

On assimile un skieur à un point matériel S de masse $m = 70$ kg. Il aborde une piste circulaire de rayon $R = 60$ m sans vitesse initiale à partir d'un angle $\alpha_0 = 25^\circ$ par rapport à l'horizontale. Il quitte la piste pour $\alpha_1 = 110^\circ$. On suppose que la piste est suffisamment verglacé pour négliger tout frottement sur la piste.

1. Établir l'équation du mouvement du skieur lorsqu'il est sur la piste. On pourra appliquer le théorème du moment cinétique.
2. En déduire l'expression de la vitesse du skieur en fonction de l'angle α repérant le skieur par rapport à l'horizontale.
3. Calculer la vitesse maximale atteinte par le skieur.

B. Exercices et problème

1. Étude du mouvement radial dans un potentiel central

Un point matériel M de masse m se déplace dans un champ de forces dérivant d'un potentiel $E_p(r) = -\frac{Km}{r}$ où K est une constante positive et $r = OM$ (O centre du repère). On note σ_O le moment cinétique de M par rapport à O .

1. Calculer l'énergie mécanique E_m et faire apparaître l'énergie potentielle efficace $E_{p,\text{eff}}$ que l'on exprimera en fonction de r , σ_O , m et K . Montrer que le mouvement ne peut avoir lieu que si $E_m \geq E_{p,\text{eff}}$

- À l'aide de l'inégalité précédente, montrer que si l'énergie mécanique est négative, la distance r du point matériel au centre attractif O reste comprise entre deux valeurs r_1 et r_2 à déterminer.
- Montrer que dans ces conditions, σ_O ne peut être supérieur à une valeur σ_m que l'on calculera. Quelle est la trajectoire pour $\sigma_O = \sigma_m$?
- Montrer que si $E_m = 0$, on a $r > r_m$ où r_m est à calculer en fonction de K , m et σ_0 .

2. Pendule sur un plan incliné

Un pendule simple est constitué d'un point M de masse m attaché à un fil de masse négligeable, de longueur L . L'autre extrémité du fil est accrochée à un point O fixe. L'ensemble peut se déplacer sans frottement sur un plan incliné faisant un angle α avec le plan horizontal (figure 25.4). On pourra introduire une base polaire $(\vec{u}_r, \vec{u}_\theta)$ non représentée, avec $\vec{u}_r = \vec{OM}/OM$ et $\vec{u}_\theta$ le vecteur directement perpendiculaire dans le plan perpendiculaire à Oz .

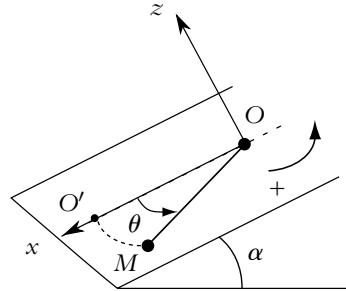


Figure 25.4

- Exprimer le moment cinétique de M par rapport à O .
- a) Déterminer la projection des différentes forces sur la base $(\vec{u}_r, \vec{u}_\theta, \vec{u}_z)$.
b) En appliquant le théorème du moment cinétique en O , déterminer l'équation différentielle du mouvement et la période des petites oscillations. Effectuer l'application numérique pour $\alpha = 30^\circ$.
c) On lance depuis $\theta = 0$ le pendule avec une vitesse v_0 . Déterminer l'angle maximal atteint, en supposant qu'on reste dans le domaine des petites oscillations.

3. Ressort en rotation d'après ENSTIM 2002

Le mouvement est étudié dans le référentiel du laboratoire assimilé à un référentiel galiléen et associé à un repère $(O, \vec{u}_x, \vec{u}_y, \vec{u}_z)$. Un palet M de masse m peut se mouvoir sans frottement dans le plan (O, x, y, z) horizontal (table à coussin d'air par exemple). Le champ de pesanteur est suivant la verticale Oz : $\vec{g} = -g\vec{u}_z$.

Un point M (masse m) est fixé à un ressort de raideur k et de longueur à vide ℓ_0 , l'autre extrémité étant fixée en O .

- Faire un bilan des forces. Montrer qu'il y a conservation du moment cinétique $\vec{L}_0$ par rapport à O .
- À $t = 0$, la masse est lâchée, sans vitesse initiale d'une longueur :

$$\ell_i = 1, 2\ell_0 \quad \text{soit} \quad \vec{OM}(t=0) = 1, 2\ell_0 \vec{u}_x$$

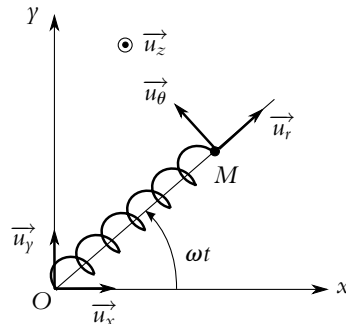


Figure 25.5

- a) Calculer $\vec{L}_0$. Quelle est la nature de la trajectoire ?
- b) Déterminer l'évolution temporelle de la longueur du ressort, $\ell(t) = OM(t)$. Préciser l'intervalle de variation de ℓ .
3. On lance la particule d'un point $\vec{OM}_0 = \vec{OM}(t=0) = \ell_1 \vec{u}_x$, avec une vitesse initiale $\vec{v}_0 = \ell_1 \omega \vec{u}_y$, orthogonale à $\vec{OM}_0$. Dans la suite, on travaillera en coordonnées polaires dans le plan (O, x, y) .
- a) Préciser $\vec{L}_0$ en fonction de $r, \dot{\theta}$ puis en fonction des conditions initiales et des vecteurs de base. On notera L le module de $\vec{L}_0$.
- b) Rappeler l'expression de l'énergie potentielle élastique. Doit-on tenir compte de l'énergie potentielle de pesanteur pour étudier le mouvement ?
Montrer qu'il y a conservation de l'énergie mécanique E_m .
Préciser l'expression de E_m :
- en fonction des conditions initiales,
 - en fonction de $r, \dot{r}, \dot{\theta}, m, k$ et ℓ_0 .
- c) Montrer que l'énergie mécanique peut s'écrire : $E_m = \frac{1}{2} m \dot{r}^2 + E_{\text{eff}}(r)$.
Préciser l'expression de $E_{\text{eff}}(r)$ et tracer son allure.
- d) La masse peut-elle s'éloigner indéfiniment du pôle d'attraction ?
- e) La vitesse de la particule peut-elle s'annuler au cours du mouvement ?
- f) La particule peut-elle passer par le centre d'attraction au cours de son mouvement ?
4. On cherche à déterminer une condition entre ℓ_1 et ω pour avoir un mouvement circulaire.
- a) Montrer que dans ce cas, le mouvement est uniforme.
- b) Déterminer ℓ_1 en fonction de k, ℓ_0 et ω . Est-elle valable pour tout ω ?

26

Changement de référentiel : aspects cinématiques

D'après la définition même du vecteur vitesse : $\vec{v} = \frac{d\vec{OM}}{dt}$ où O est un point fixe dans le référentiel dans lequel on la détermine, il dépend du référentiel par rapport auquel on la calcule. De même, l'accélération : $\vec{a} = \frac{d^2\vec{OM}}{dt^2}$ dépend du référentiel.

L'objet de ce chapitre est d'établir le lien entre les vitesses et les accélérations exprimées dans deux référentiels différents, en se limitant aux aspects cinématiques.

1. Position du problème

1.1 Exemples

Les changements de référentiel sont une question qui apparaît naturellement dans des situations de la vie courante.

Ainsi l'étude de la chute d'un objet dans un train en translation uniforme peut être considérée soit par rapport à un observateur assis dans le train, soit par rapport à un observateur placé le long de la voie de chemin de fer. Les deux observateurs constituent deux référentiels possibles. Chacun d'eux ne verra pas la même trajectoire. On peut établir que l'observateur assis dans le train voit l'objet tomber le long d'une verticale alors que celui placé le long de la voie verra une trajectoire parabolique.

De même, si on considère le mouvement d'une valve de roue de voiture, la trajectoire diffère suivant le référentiel envisagé. Dans celui qui est lié au centre de la roue et qui se déplace avec le véhicule, il s'agit d'un cercle. Dans celui qui est lié au sol, on montre que la valve décrit une cycloïde lorsque la voiture ne glisse pas sur le sol :



Figure 26.1 Trajectoire d'une valve de roue en mouvement dans le référentiel lié au sol.

1.2 Formalisation du problème

L'objet du chapitre est d'établir d'un point de vue strictement cinématique (en excluant tout aspect dynamique) les lois permettant d'obtenir les expressions de la vitesse et de l'accélération dans un référentiel connaissant leur expression dans un autre référentiel et le mouvement du nouveau référentiel par rapport à l'ancien. Pour l'exemple de la valve, cela revient, par exemple, à déterminer sa vitesse et son accélération par rapport au sol connaissant celles par rapport au centre de la roue, et le mouvement du référentiel lié au centre de la roue qui se déplace avec le véhicule.

On formalise ce problème en définissant les notions :

- *de référentiel absolu :*

Le référentiel absolu est le référentiel considéré comme fixe. On le note $\mathcal{R}$ et on lui associe un observateur fixe O et un système d'axes $\vec{u}_x, \vec{u}_y, \vec{u}_z$. On notera x, y, z les composantes du point M dans ce référentiel.

- *de référentiel relatif :*

Le référentiel relatif est le référentiel considéré comme mobile. On le note $\mathcal{R}'$ et on lui associe un observateur fixe O' dans ce référentiel et un système d'axes $\vec{u}'_x, \vec{u}'_y, \vec{u}'_z$. On notera x', y', z' les composantes du point M dans ce référentiel.

Ces dénominations sont formelles et n'ont pas de signification intrinsèque. Les deux référentiels ont des rôles parfaitement symétriques, le seul point important est l'existence d'un mouvement relatif de l'un par rapport à l'autre. Par exemple, dans le cas de la chute d'un objet dans un train, on peut considérer que le référentiel absolu est celui lié à l'observateur le long de la voie et que le train se déplace par rapport à ce dernier. Une autre manière d'envisager le problème consiste à prendre le train fixe et à considérer que l'observateur le long de la voie se déplace. Dans ce cas, le référentiel lié au train est le référentiel absolu et celui lié à l'observateur le long de la voie est le référentiel relatif.

1.3 Notion de point coïncident

On appelle point coïncident le point P du référentiel mobile qui, à l'instant t , coïncide avec le point M du référentiel fixe.

Le point P est fixe dans le référentiel mobile : il lui est rigidement lié. En revanche, le point M défini dans le référentiel fixe est mobile : il n'est pas au même endroit pour deux instants distincts.

Par exemple, dans le cas d'une roue qui tourne, le point P est le point appartenant à la roue qui est fixe par rapport à elle et le point M est le point du sol qui se confond avec P à l'instant où P est au sol.

— à l'instant t
 à l'instant $t + dt$

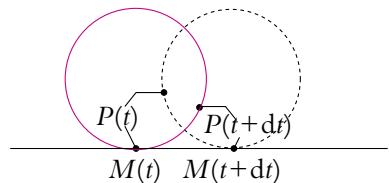


Figure 26.2 Point coïncident d'une roue.

Il est à noter qu'en général si P et M coïncident à l'instant t , il n'y a aucune raison qu'ils coïncident à un instant antérieur ou postérieur. Leurs trajectoires se croisent à l'instant t et se séparent juste après.

Le point P est entraîné par le référentiel mobile auquel il est lié tandis que le point M reste fixe. P coïncide avec un autre point M' à l'instant $t+dt$.

1.4 Lien avec le torseur cinématique d'entraînement vu en Sciences Industrielles

On peut lier le problème des changements de référentiel à l'étude des mouvements des solides effectuée dans le cours de sciences industrielles et notamment à la notion de torseur cinématique d'entraînement.

En effet, le mouvement d'un référentiel par rapport à un autre peut être considéré comme celui d'un solide : l'observateur étant l'analogue d'un point particulier du solide et les axes du référentiel définissant l'orientation du solide.

Or le mouvement d'un solide est décrit par un torseur cinématique à savoir :

- la vitesse de l'un quelconque de ses points,
- le vecteur rotation du solide.

On peut donc décrire le mouvement d'un référentiel par rapport à un autre par un torseur cinématique d'entraînement à savoir :

- la vitesse d'un point fixe du référentiel mobile par rapport au référentiel fixe,
- le vecteur rotation du référentiel mobile par rapport au référentiel fixe.

On verra que ces deux quantités suffisent à décrire le mouvement du référentiel mobile par rapport au référentiel fixe et permettent une description complète du passage du référentiel fixe au référentiel mobile.

1.5 Définitions

On distingue plusieurs vitesses et accélérations.

- La vitesse (respectivement l'accélération) absolue notée $\vec{v}_{/\mathcal{R}}(M)$ (respectivement $\vec{a}_{/\mathcal{R}}(M)$) est la vitesse (respectivement l'accélération) du point M étudié dans le référentiel fixe ou absolu $\mathcal{R}$.
- La vitesse (respectivement l'accélération) relative notée $\vec{v}_{/\mathcal{R}'}(M)$ (respectivement $\vec{a}_{/\mathcal{R}'}(M)$) est la vitesse (respectivement l'accélération) du point M étudié dans le référentiel mobile ou relatif $\mathcal{R}'$.
- La vitesse (respectivement l'accélération) d'entraînement notée $\vec{v}_e(M)$ (respectivement $\vec{a}_e(M)$) est la vitesse (respectivement l'accélération) du point coïncidant P évaluée dans le référentiel fixe ou absolu $\mathcal{R}$.



On verra que l'accélération d'entraînement n'est pas égale à la dérivée de la vitesse d'entraînement.

2. Composition des vitesses

Il s'agit maintenant d'établir le lien entre ces différentes quantités. On commence par la relation liant les différentes vitesses.

Soit $\mathcal{R}$ lié à un observateur O et à une base $(\vec{u}_x, \vec{u}_y, \vec{u}_z)$ fixes.

Soit $\mathcal{R}'$ lié à un observateur O' et à une base $(\vec{u}'_x, \vec{u}'_y, \vec{u}'_z)$ mobiles par rapport à $\mathcal{R}$.

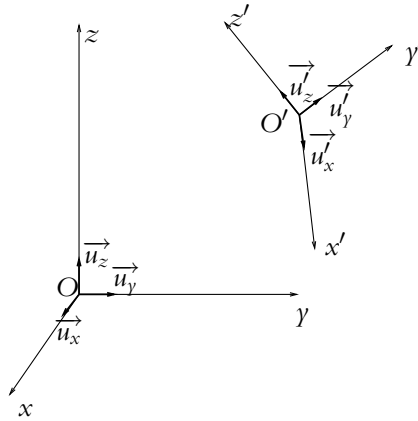


Figure 26.3 Changement de référentiel.

2.1 Loi de composition des vitesses

Dans le référentiel fixe $\mathcal{R}$, le vecteur-position s'écrit :

$$\overrightarrow{OM} = x(t)\vec{u}_x + y(t)\vec{u}_y + z(t)\vec{u}_z$$

On en déduit le vecteur-vitesse du point M par rapport au référentiel $\mathcal{R}$:

$$\vec{v}_{/\mathcal{R}}(M) = \dot{x}\vec{u}_x + \dot{y}\vec{u}_y + \dot{z}\vec{u}_z$$

On procède de même dans le référentiel mobile $\mathcal{R}'$ pour obtenir l'expression du vecteur-vitesse du point M par rapport au référentiel $\mathcal{R}'$:

$$\overrightarrow{O'M} = x'(t)\vec{u}'_x + y'(t)\vec{u}'_y + z'(t)\vec{u}'_z$$

et

$$\vec{v}_{/\mathcal{R}'}(M) = \dot{x}'\vec{u}'_x + \dot{y}'\vec{u}'_y + \dot{z}'\vec{u}'_z$$

Or

$$\overrightarrow{OM} = \overrightarrow{OO'} + \overrightarrow{O'M} = \overrightarrow{OO'} + x'\vec{u}'_x + y'\vec{u}'_y + z'\vec{u}'_z$$

Alors en dérivant cette expression :

$$\begin{aligned}
 \vec{v}_{/\mathcal{R}}(M) &= \left(\frac{d\overrightarrow{OM}}{dt} \right)_{\mathcal{R}} \\
 &= \left(\frac{d\overrightarrow{OO'}}{dt} \right)_{\mathcal{R}} + x' \vec{u}_x + x' \frac{d\vec{u}_x}{dt} + y' \vec{u}_y + y' \frac{d\vec{u}_y}{dt} + z' \vec{u}_z + z' \frac{d\vec{u}_z}{dt} \\
 &= \left(\frac{d\overrightarrow{OO'}}{dt} \right)_{\mathcal{R}} + x' \frac{d\vec{u}_x}{dt} + y' \frac{d\vec{u}_y}{dt} + z' \frac{d\vec{u}_z}{dt} + \vec{v}_{/\mathcal{R}'}(M)
 \end{aligned}$$

Soit P le point coïncidant, alors $\overrightarrow{OP} = x' \vec{u}_x + y' \vec{u}_y + z' \vec{u}_z$ avec x' , y' et z' constantes dans le référentiel $\mathcal{R}'$ par définition même de ce point. On a donc : $\vec{v}_{/\mathcal{R}'}(P) = \vec{0}$. En appliquant la relation précédente au point coïncidant, on obtient l'expression de sa vitesse du point évaluée dans le référentiel fixe $\mathcal{R}$:

$$\vec{v}_{/\mathcal{R}}(P) = \left(\frac{d\overrightarrow{OO'}}{dt} \right)_{\mathcal{R}} + x' \frac{d\vec{u}_x}{dt} + y' \frac{d\vec{u}_y}{dt} + z' \frac{d\vec{u}_z}{dt}$$

Il s'agit par définition de la *vitesse d'entraînement*.

On en déduit la relation de la loi de composition des vitesses :

$$\vec{v}_{/\mathcal{R}}(M) = \vec{v}_{/\mathcal{R}'}(M) + \vec{v}_e(M)$$

2.2 Cas où $\mathcal{R}'$ est en translation par rapport à $\mathcal{R}$

Dans ce cas, les axes ne subissent aucune rotation donc les vecteurs unitaires sont fixes dans le référentiel $\mathcal{R}$:

$$\frac{d\vec{u}_x}{dt} = \frac{d\vec{u}_y}{dt} = \frac{d\vec{u}_z}{dt} = \vec{0}$$

et on en déduit l'expression de la vitesse d'entraînement :

$$\vec{v}_e(M) = \left(\frac{d\overrightarrow{OO'}}{dt} \right)_{\mathcal{R}}$$

Elle est alors indépendante du point M .

La vitesse d'entraînement n'est dans ce cas que la vitesse de l'origine O' (ou de tout autre point fixe par rapport au référentiel mobile) par rapport au référentiel fixe.

2.3 Cas où $\mathcal{R}'$ est en rotation autour d'un axe fixe de $\mathcal{R}$

Sans perte de généralité du problème, on peut choisir O' confondu avec O : il suffit de prendre un point fixe dans les deux référentiels, ce qui est toujours possible dans le cas d'une rotation autour d'un axe fixe puisque tous les points de l'axe de rotation le sont.

On choisit en outre l'axe Oz selon l'axe de rotation de $\mathcal{R}'$ par rapport à $\mathcal{R}$.

On note θ l'angle entre les axes (Ox) et (Ox') à l'instant t et $\omega(t) = \frac{d\theta}{dt}$. On remarque que $\omega(t)$ n'est pas forcément indépendant du temps : il suffit que la rotation n'ait pas lieu à vitesse angulaire constante. On peut traduire la rotation du référentiel mobile en introduisant le vecteur rotation

$$\vec{\Omega} = \omega \vec{u}_z$$

Par définition, la vitesse d'entraînement est la vitesse du point coïncidant. Or dans le cas où $\mathcal{R}'$ est en rotation autour d'un axe fixe de $\mathcal{R}$, le point coïncidant P décrit un mouvement circulaire à la vitesse angulaire $\omega(t)$.

On en déduit l'expression de la vitesse d'entraînement :

$$\vec{v}_e(M) = \vec{\Omega} \wedge \overrightarrow{O'M}$$

La notion de point coïncidant est donc particulièrement efficace ici pour déterminer l'expression de la vitesse d'entraînement

2.4 Cas général

On peut démontrer que le mouvement le plus général d'un référentiel $\mathcal{R}'$ par rapport à un référentiel fixe $\mathcal{R}$ peut être décomposé à tout instant en :

- un mouvement de rotation instantanée autour d'un axe,
- un mouvement de translation dans la direction définie par cet axe autour duquel s'effectue la rotation instantanée.

Du fait de cette décomposition, on obtient l'expression vectorielle intrinsèque de la vitesse d'entraînement :

$$\vec{v}_e(M) = \left(\frac{d\overrightarrow{OO'}}{dt} \right)_{\mathcal{R}} + \vec{\Omega} \wedge \overrightarrow{O'M}$$

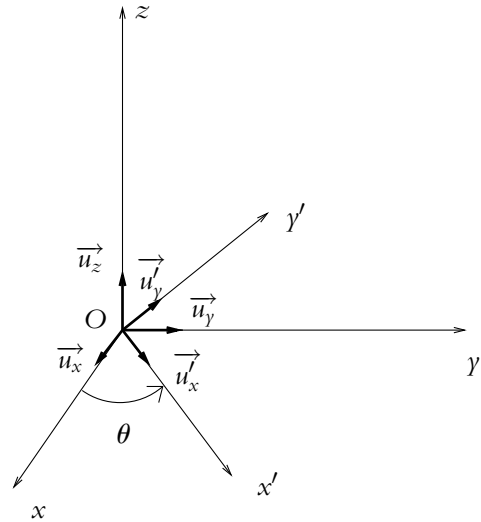


Figure 26.4 Changement de référentiel par rotation.

3. Composition des accélérations

3.1 Loi de composition des accélérations

On reprend les expressions obtenues au paragraphe précédent pour le vecteur-vitesse de M par rapport au référentiel $\mathcal{R}$:

$$\vec{v}_{/\mathcal{R}}(M) = \dot{x}\vec{u}_x + \dot{y}\vec{u}_y + \dot{z}\vec{u}_z$$

$$\vec{v}_{/\mathcal{R}}(M) = \frac{d\vec{OO'}}{dt} + x'\frac{d\vec{u}'_x}{dt} + y'\frac{d\vec{u}'_y}{dt} + z'\frac{d\vec{u}'_z}{dt} + \dot{x}'\vec{u}'_x + \dot{y}'\vec{u}'_y + \dot{z}'\vec{u}'_z \quad (26.1)$$

qu'on dérive par rapport au temps pour obtenir l'accélération de M par rapport au référentiel $\mathcal{R}$:

$$\begin{aligned} \vec{a}_{/\mathcal{R}}(M) = & \frac{d^2\vec{OO'}}{dt^2} + x'\frac{d^2\vec{u}'_x}{dt^2} + y'\frac{d^2\vec{u}'_y}{dt^2} + z'\frac{d^2\vec{u}'_z}{dt^2} + 2\dot{x}'\frac{d\vec{u}'_x}{dt} + 2\dot{y}'\frac{d\vec{u}'_y}{dt} \\ & + 2\dot{z}'\frac{d\vec{u}'_z}{dt} + \ddot{x}'\vec{u}'_x + \ddot{y}'\vec{u}'_y + \ddot{z}'\vec{u}'_z \end{aligned}$$

On identifie

- le terme d'accélération de M par rapport au référentiel $\mathcal{R}'$:

$$\vec{a}_{/\mathcal{R}'}(M) = \ddot{x}'\vec{u}'_x + \ddot{y}'\vec{u}'_y + \ddot{z}'\vec{u}'_z$$

- un terme dit d'*accélération d'entraînement* correspondant à l'accélération du point coïncidant obtenu en annulant les dérivées premières et secondes de x' , y' et z' :

$$\vec{a}_e(M) = \frac{d^2\vec{OO'}}{dt^2} + x'\frac{d^2\vec{u}'_x}{dt^2} + y'\frac{d^2\vec{u}'_y}{dt^2} + z'\frac{d^2\vec{u}'_z}{dt^2}$$

- un dernier terme appelé *accélération complémentaire ou de Coriolis* :

$$\vec{a}_c(M) = 2\dot{x}'\frac{d\vec{u}'_x}{dt} + 2\dot{y}'\frac{d\vec{u}'_y}{dt} + 2\dot{z}'\frac{d\vec{u}'_z}{dt}$$

La loi de composition des accélérations s'exprime par la relation :

$$\vec{a}_{/\mathcal{R}}(M) = \vec{a}_{/\mathcal{R}'}(M) + \vec{a}_e(M) + \vec{a}_c(M)$$



Attention, en général, l'accélération d'entraînement n'est pas la dérivée par rapport au temps de la vitesse d'entraînement.

On cherche à donner une interprétation physique intrinsèque des accélérations d'entraînement et de Coriolis. On commence pour ce faire par étudier les cas particuliers où le référentiel $\mathcal{R}'$ est en translation par rapport au référentiel $\mathcal{R}$, puis celui où le référentiel $\mathcal{R}'$ est en rotation autour d'un axe fixe de $\mathcal{R}$.

3.2 Cas où $\mathcal{R}'$ est en translation par rapport à $\mathcal{R}$

Dans le cas d'une translation de $\mathcal{R}'$ par rapport à $\mathcal{R}$, $(\vec{u}'_x, \vec{u}'_y, \vec{u}'_z)$ gardent des directions fixes par rapport à $(\vec{u}_x, \vec{u}_y, \vec{u}_z)$ donc les dérivées de ces vecteurs par rapport au temps sont nulles.

On en déduit les expressions des accélérations d'entraînement et de Coriolis dans ce cas :

$$\begin{cases} \vec{a}_e(M) = \frac{d^2 \vec{OO'}}{dt^2} \\ \vec{a}_c(M) = \vec{0} \end{cases}$$

L'accélération d'entraînement n'est autre que l'accélération de l'origine ou de n'importe quel point fixe dans le référentiel mobile par rapport au référentiel fixe.

L'accélération de Coriolis est nulle dans ce cas.

3.3 Cas où $\mathcal{R}'$ est en rotation par rapport à $\mathcal{R}$

Comme pour l'étude de la composition des vitesses, on se place dans le cadre suivant (qui ne réduit pas le caractère général du résultat) :

- on utilise les coordonnées cartésiennes,
- on choisit l'axe Oz selon l'axe de rotation,
- on note θ l'angle de rotation et $\vec{\Omega} = \omega \vec{u}'_z = \frac{d\theta}{dt} \vec{u}'_z$ le vecteur rotation instantané.

On va établir à partir de ces hypothèses une expression intrinsèque des accélérations d'entraînement et de Coriolis.

a) Accélération d'entraînement

L'accélération d'entraînement est celle du point coïncidant qui, dans le cas où $\mathcal{R}'$ est en rotation par rapport à $\mathcal{R}$, a une trajectoire circulaire de rayon $r = HM$, en notant H la projection orthogonale de M sur l'axe de rotation Oz , sa vitesse angulaire étant $\omega(t)$. L'accélération d'entraînement est donc celle d'un mouvement circulaire, à savoir en coordonnées polaires dans le plan perpendiculaire à Oz passant par M :

$$\vec{a}_e(M) = -r\omega^2 \vec{u}_r + r \frac{d\omega}{dt} \vec{u}_\theta$$

qu'on peut écrire sous la forme :

$$\vec{a}_e(M) = -\omega^2 \vec{HM} + \frac{d\vec{\Omega}}{dt} \wedge \vec{O'M}$$

car $\vec{O'M} = \vec{O'H} + \vec{HM}$ avec $\vec{O'H}$ parallèle à l'axe Oz donc à $\vec{\Omega} = \omega(t)\vec{u}_z$, et

$$\frac{d\vec{\Omega}}{dt} \wedge \vec{O'M} = \frac{d\vec{\Omega}}{dt} \wedge \vec{HM} = \frac{d\omega}{dt} \vec{u}_z \wedge r\vec{u}_r = \frac{d\omega}{dt} r\vec{u}_\theta$$

On remarque que l'accélération d'entraînement n'est pas la dérivée par rapport au temps de la vitesse d'entraînement.

On retiendra l'expression de l'accélération d'entraînement

$$\vec{a}_e(M) = \frac{d\vec{\Omega}}{dt} \wedge \vec{O'M} - \Omega^2 \vec{HM}$$

en notant H la projection de M sur l'axe Oz de rotation de $\mathcal{R}'$ par rapport à $\mathcal{R}$.

On peut également écrire :

$$\vec{a}_e(M) = \frac{d\vec{\Omega}}{dt} \wedge \vec{O'M} + \vec{\Omega} \wedge (\vec{\Omega} \wedge \vec{O'M}).$$

On note la présence de deux termes dont l'un dépend des variations du vecteur rotation au cours du temps.



Dans le cas particulier où la rotation est uniforme :

$$\vec{a}_e(M) = -\Omega^2 \vec{HM}$$

b) Accélération de Coriolis

On utilise les résultats sur la dérivation d'un vecteur unitaire par rapport à un angle

de rotation : $\frac{d\vec{u}'_x}{dt} = \omega\vec{u}'_y$, $\frac{d\vec{u}'_y}{dt} = -\omega\vec{u}'_x$ et $\frac{d\vec{u}'_z}{dt} = \vec{0}$,

l'accélération de Coriolis s'écrit alors :

$$\vec{a}_c(M) = 2\dot{x}'\omega\vec{u}'_y - 2\dot{y}'\omega\vec{u}'_x$$

Or la vitesse de M par rapport au référentiel $\mathcal{R}'$ a pour expression :

$$\vec{v}_{/\mathcal{R}'}(M) = \dot{x}'\vec{u}'_x + \dot{y}'\vec{u}'_y + \dot{z}'\vec{u}'_z$$

donc

$$2\vec{\Omega} \wedge \overrightarrow{v_{/\mathcal{R}'}}(M) = 2\omega u'_z \wedge \overrightarrow{v_{/\mathcal{R}'}}(M) = 2x'\omega u'_y - 2y'\omega u'_x$$

Le terme de Coriolis

$$\vec{a}_c(M) = 2\vec{\Omega} \wedge \overrightarrow{v_{/\mathcal{R}'}}(M)$$

est donc lié au mouvement relatif d'un point dans le référentiel mobile. Pour un point fixe, ce terme est nul.

3.4 Cas général

La décomposition de tout mouvement sous la forme d'un mouvement de rotation autour d'un axe appelé axe instantané de rotation et de glissement et d'un mouvement de translation le long de cet axe conduit aux expressions générales pour les accélérations d'entraînement :

$$\vec{a}_e(M) = \left(\frac{d^2 \overrightarrow{OO'}}{dt^2} \right)_{\mathcal{R}} + \frac{d\vec{\Omega}}{dt} \wedge \overrightarrow{O'M} + \vec{\Omega} \wedge (\vec{\Omega} \wedge \overrightarrow{O'M})$$

et de Coriolis :

$$\vec{a}_c(M) = 2\vec{\Omega} \wedge \overrightarrow{v_{/\mathcal{R}'}}(M)$$



L'expression de l'accélération d'entraînement n'est **que rarement utile sous sa forme générale** : il faut savoir que c'est l'accélération du point coïncidant et il suffit de la connaître dans les deux cas simples étudiés précédemment : le référentiel $\mathcal{R}'$ est en translation par rapport au référentiel $\mathcal{R}$, ou en rotation autour d'un axe fixe de $\mathcal{R}$.

En revanche, **il faut connaître** l'expression générale de l'accélération de Coriolis.

A. Applications directes du cours

1. Mesure de vitesse à l'aide d'observations

Le passager d'une voiture observe que la neige tombe en formant un angle de 10° par rapport à la verticale lorsque celui-ci roule à 110 km.h^{-1} . Lorsque la voiture s'arrête au feu, le passager regarde la neige tomber et constate que celle-ci tombe verticalement.

Calculer la vitesse de la neige par rapport au sol puis par rapport à la voiture lorsqu'il roule.

2. Traversée d'un tapis roulant

Lors d'un jeu télévisé, un joueur A doit traverser un tapis roulant de largeur a , pour donner un paquet à un second joueur B . Le tapis se déplace à une vitesse constante V_t par rapport au sol. Lorsque le joueur court sur le tapis, sa vitesse par rapport au tapis a pour norme V constante.

1. Le joueur A se déplace avec une vitesse $\vec{V}$ perpendiculaire au bord du tapis. Où doit se placer B pour réceptionner le paquet ? Quel est le temps t_1 de traversée du tapis ?
2. Pour le deuxième essai, le joueur B est posté en face du joueur A . Dans quelle direction A doit-il courir ? Quel est le temps de traversée t_2 ?
3. On suppose maintenant que la vitesse $\vec{V}$ fait un angle θ quelconque avec $\vec{V}_t$. Déterminer le temps de traversée t_3 en fonction de a , V et θ . Montrer que c'est t_1 le temps de traversée le plus court.

3. Manège et spirale d'Archimède

Un enfant se déplace le long d'un diamètre AB d'un manège circulaire tournant à vitesse angulaire ω constante. Sa vitesse par rapport au manège est constante.

1. Déterminer la trajectoire de l'enfant dans le référentiel du manège.
2. Même question dans le référentiel lié au sol.
3. L'enfant traverse le manège en un tour de manège. Tracer la courbe qu'il décrit par rapport au sol.
4. Calculer dans le référentiel lié au sol la vitesse et l'accélération de l'enfant.

4. Double rotation

Une tige est en rotation autour de son centre O . Un point matériel M est en rotation autour de A , une des extrémités de la tige. On suppose que les deux rotations ont lieu dans un même plan et que la vitesse de rotation de M autour de A est plus grande que celle de la tige autour de son centre. Déterminer l'allure des trajectoires de A et de M .

5. Roulement sans glissement d'une roue

Un point O se déplace à vitesse constante v_0 suivant Ox dans un référentiel fixe. Un point M décrit un cercle de centre O et de rayon R à la vitesse angulaire ω_0 . Quelle est la trajectoire de M dans le référentiel fixe ? On distinguera trois cas :

1. $v_0 < R\omega_0$,
2. $v_0 = R\omega_0$,
3. $v_0 > R\omega_0$.

B. Exercices et problèmes

1. Sauvetage d'un véliplanchiste

Un véliplanchiste imprudent s'est aventuré trop loin au large. Il est animé d'un mouvement rectiligne uniforme de vitesse $\vec{v}$. Les secouristes sur la plage montent dans leur ZODIAC pour aller le chercher. On note N la position de la planche à voile et P celle du poste de sauvetage. A l'instant où l'angle entre la vitesse $\vec{v}$ et $\overrightarrow{NP}$ prend la valeur α , les secouristes démarrent. La vitesse u de leur bateau est supposée constante.

1. Quelle doit être la valeur de l'angle avec lequel se déplace le bateau pour qu'il atteigne la planche à voile ?
2. Quelle sera la valeur choisie si on veut que les secouristes atteignent leur but le plus vite possible ?

2. Mouvement d'un bac sur un fleuve

Un bac se déplace à la vitesse $\vec{v}$ par rapport à l'eau d'un fleuve qui coule à la vitesse $\vec{u}$ par rapport aux rives. On suppose que le module v de la vitesse est constant et indépendant de sa direction. Il met une durée T_1 pour traverser le fleuve de largeur d , son mouvement étant perpendiculaire aux rives dans le référentiel lié au sol. Pour parcourir la même distance parallèlement au sens du courant, il met une durée T_2 pour descendre dans le sens du courant et une durée T_3 pour remonter le courant.

1. Exprimer $\frac{T_1}{T_2}$ et $\frac{T_3}{T_1}$ en fonction de u et v .
2. Déterminer la vitesse du bac par rapport à l'eau en fonction de u sachant que $T_1 = 3 T_2$.
3. Déterminer, dans chacun des cas, les vitesses du bac par rapport au rivage en fonction de u .
4. Que vaut alors $\frac{T_3}{T_1}$? Conclure à partir des valeurs respectives des vitesses absolues correspondant aux durées T_2 et T_3 .

3. Arme de jet

Un trébuchet (figure 26.5) était une arme du Moyen-Âge qui permettait d'envoyer des charges lourdes contre les murailles. Il est composé d'une poutre BC (appelée verge) à laquelle est fixée un contrepoids en B (la huche). En C est attachée une corde au bout de laquelle une poche contient le projectile M (boulet de pierre taillée).

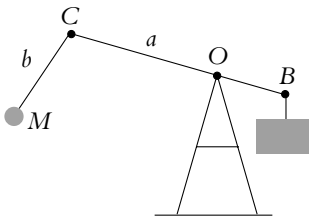


Figure 26.5

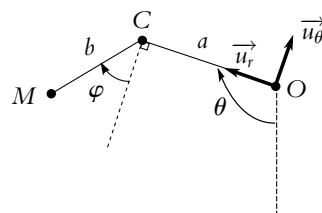


Figure 26.6

Les différents paramètres sont définis sur la figure 26.6. En particulier, la direction de OC est repérée par θ par rapport à la verticale, et la direction de CM par φ par rapport à la direction perpendiculaire à OC . On note $\mathcal{R}_S$ le référentiel fixe lié au sol et $\mathcal{R}_C$ celui lié au point C et à la base $(\vec{u}_r, \vec{u}_\theta)$.

1. Quel est le mouvement de $\mathcal{R}_C$ par rapport à $\mathcal{R}_S$?
2. On suppose que la corde CM reste tendue. En déduire la vitesse de M dans $\mathcal{R}_C$ projetée sur $(\vec{u}_r, \vec{u}_\theta)$ en fonction de b , $\dot{\varphi}$ et φ .
3. Déterminer le vecteur $\vec{OM}$ en projection sur $(\vec{u}_r, \vec{u}_\theta)$. En déduire la vitesse d'entraînement en M pour le mouvement de $\mathcal{R}_C$ par rapport à $\mathcal{R}_S$.
4. Le projectile est lâché lorsque $\theta = \pi$ et $\varphi = \pi/2$ ($BOCM$ vertical). Déterminer alors la vitesse de M dans $\mathcal{R}_S$ en fonction de a , b , $\dot{\varphi}$ et $\dot{\theta}$. Montrer que la vitesse obtenue est plus grande que s'il y avait un seul bras rigide de longueur $a + b$.

4. Avion confronté à un vent contraire, d'après Agrégation de Chimie 2006

On considère le sol comme un référentiel galiléen (référentiel absolu) noté $\mathcal{R}$.

Un avion doit se déplacer en ligne droite d'un point A vers un point au sol B (Figure 26.7). Il subit un vent contraire de vitesse v_e . Le vecteur $\vec{v}_e$ fait un angle φ avec la trajectoire AB . L'avion vole à une vitesse constante $\vec{V}$ par rapport à l'air. Cette vitesse fait un angle θ avec la route au sol AB .

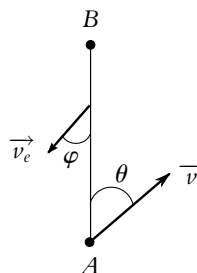


Figure 26.7

1. a) Calculer la vitesse $\vec{v}_a$ de l'avion dans le référentiel $\mathcal{R}$ en fonction des données.
- b) À quelles conditions l'avion peut-il se déplacer en ligne droite de A vers B ?
- c) Calculer l'angle de correction θ que le pilote doit afficher dans le cas où $V = 445 \text{ km.h}^{-1}$, pour contrer un vent de vitesse $v_e = 56 \text{ km.h}^{-1}$ et $\varphi = 20^\circ$.
2. L'avion doit faire un aller-retour entre les deux points A et B distants de $d = 500 \text{ km}$, dans les conditions de la question précédente.
 - a) Calculer la durée t_{ar} du trajet aller-retour. On négligera le temps du demi-tour.
 - b) Comparer cette durée avec le temps t'_{ar} qu'aurait mis l'avion pour faire le même trajet en l'absence de vent. Est-il possible que $t_{ar} < t'_{ar}$?
Commenter cette maxime de l'aéronautique : "le temps perdu ne se rattrape jamais".
3. En B , l'avion arrive au-dessus d'une balise au sol à une date prise pour origine des temps. Le contrôleur demande alors au pilote de réaliser à altitude constante un circuit d'attente $BCDE$ (Figure 26.8) constitué de deux parties semi-circulaires BC et DE et de deux parties rectilignes CD et EB . Il lui indique à quelle date t_B l'avion doit se présenter à nouveau au-dessus de la balise B .

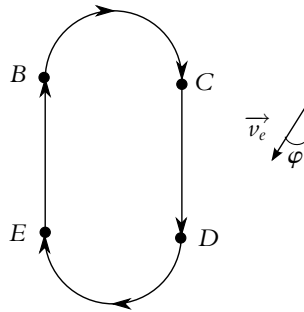


Figure 26.8

Le pilote adapte l'inclinaison de l'appareil et donc le rayon des tronçons semi-circulaires BC et DE de manière à ce que ceux-ci durent 1 minute chacun. Durant cette attente, la vitesse V de l'avion est supposée constante et égale à 222 km.h^{-1} , tandis que la vitesse du vent est toujours $v_e = 56 \text{ km.h}^{-1}$, et $\varphi = 20^\circ$.

À quelle date t_v le pilote doit-il entreprendre le virage DE si la date de passage au dessus du point B , imposée par le contrôleur, est $t_B = 4$ minutes ?

5. Engrenage

On s'intéresse à un système de deux roues dentées modélisé par deux cercles : un grand de rayon R et centre O immobile qui sert de guide au plus petit de centre C et de rayon r . Le centre C de la petite roue à un mouvement circulaire uniforme de vitesse v_0 . On s'intéresse au mouvement d'un point M de la périphérie de la petite roue. On note I le point de contact entre les deux engrenages ; à $t = 0$, I et M se trouvent sur l'axe Ox . Ultérieurement, ces points sont repérés par les angles $\varphi = (Ox, \vec{OI})$ et $\theta = (\vec{CI}, \vec{CM})$

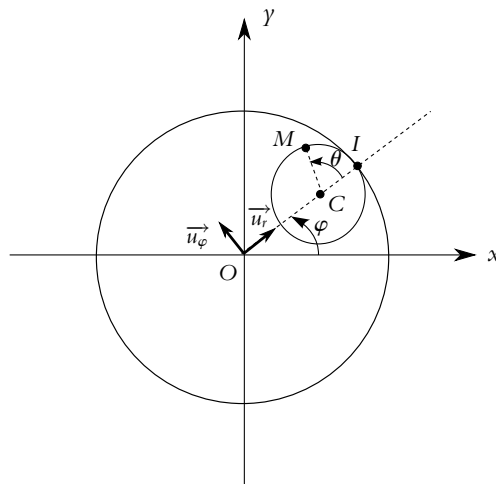


Figure 26.9

On note $\mathcal{R}$ le référentiel fixe lié à Ox, Oy et $\mathcal{R}'$ celui lié à $(C, \vec{u}_r, \vec{u}_\varphi)$.

Au point de contact, il y a roulement sans glissement, c'est-à-dire que $\vec{v}_{\mathcal{R}}(I) = \vec{0}$.

1. Déterminer la relation entre $\omega_1 = \frac{d\varphi}{dt}$, v_0 , r et R . En déduire l'expression de φ en fonction du temps.
2. Déterminer la vitesse du point M dans $\mathcal{R}'$ en projection sur $\vec{u}_r$ et $\vec{u}_\varphi$. On note $\omega_2 = \frac{d\theta}{dt}$.
3. Déterminer les composantes du vecteur $\vec{OM}$ sur $\vec{u}_r$ et $\vec{u}_\varphi$. En déduire la vitesse d'entraînement de M en projection sur ces mêmes vecteurs. Établir alors l'expression de la vitesse absolue de M : $\vec{v}_{/\mathcal{R}}(M)$.
4. Quelle relation entre ω_1 , ω_2 , R et r résulte de la condition de roulement sans glissement ?
5. Par composition des accélérations, déterminer l'expression de l'accélération absolue de M : $\vec{a}_{/\mathcal{R}}(M)$.

Changement de référentiel : aspects dynamiques

27

Ce chapitre s'intéresse aux conséquences dynamiques du problème des changements de référentiel. En première période, on a vu que la première loi de Newton ou principe d'inertie entraîne que tous les référentiels ne sont pas équivalents et qu'il existe des référentiels privilégiés : les référentiels galiléens. Ces référentiels seront donc les référentiels absolus et les référentiels non galiléens les référentiels relatifs. On va donc commencer par préciser la notion de relativité galiléenne. On pourra alors expliciter l'écriture du principe fondamental de la dynamique, du théorème du moment cinétique et de l'énergie cinétique en référentiel non galiléen.

1. Relativité galiléenne

1.1 Infinité de référentiels galiléens

La première loi de Newton ou principe d'inertie postule l'existence d'un référentiel galiléen. Il est donc naturel de chercher à prouver expérimentalement ce principe en exhibant un référentiel galiléen. Pour cela, on va tester le caractère galiléen des référentiels. On cherche à vérifier que l'application des lois de Newton dans le référentiel considéré fournit une bonne interprétation des résultats observés.

Comme on le verra au paragraphe suivant, la recherche d'un référentiel galiléen est loin d'être évidente. En revanche, on peut obtenir aisément une infinité de tels référentiels à partir du moment où on en connaît un. On va pour cela établir qu'un référentiel en translation uniforme par rapport à un référentiel galiléen est également galiléen.

Soient $\mathcal{R}$ un référentiel galiléen et $\mathcal{R}'$ un référentiel en translation uniforme à la vitesse constante $\vec{V}$ par rapport à $\mathcal{R}$.

Un point matériel isolé est animé, du fait du principe d'inertie, d'un mouvement rectiligne uniforme avec une vitesse constante notée $\vec{v}$ par rapport au référentiel

galiléen $\mathcal{R}$. La loi de composition des vitesses permet d'obtenir la vitesse $\vec{v}'$ de ce point matériel isolé dans le référentiel $\mathcal{R}'$:

$$\vec{v}_{/\mathcal{R}}(M) = \vec{v} = \vec{v}_{/\mathcal{R}'}(M) + \vec{v}_e(M) = \vec{v}' + \vec{V}$$

soit

$$\vec{v}' = \vec{v} - \vec{V}$$

Or $\vec{v}$ et $\vec{V}$ sont constantes : on en déduit que $\vec{v}'$ est une constante. Le mouvement d'un point matériel isolé est rectiligne uniforme dans le référentiel $\mathcal{R}'$. Ceci prouve le caractère galiléen de ce référentiel.

Tout référentiel en translation rectiligne uniforme par rapport à un référentiel galiléen est galiléen.

En mécanique classique, le passage d'un référentiel galiléen à un autre s'obtient donc à partir de la transformation de Galilée¹ :

$$\begin{cases} \vec{OM} = \vec{O'M} + \vec{V}t \\ t = t' \end{cases}$$

1.2 Recherche d'un référentiel galiléen

Il reste maintenant à déterminer un référentiel galiléen. La première idée qui vient à l'esprit est de s'intéresser au *référentiel terrestre* qui est défini comme un référentiel lié à la Terre. Son origine est donc un point de la Terre et ses axes ont des directions fixes par rapport à cette planète. Dans un tel référentiel, l'étude, par exemple, de la chute d'un objet conduit par application des lois de Newton² à un mouvement rectiligne uniformément accéléré le long de la verticale. La réalisation de l'expérience semble confirmer ce résultat. On serait donc tenté de dire qu'on vient de déterminer le référentiel galiléen cherché.

Cependant, comme on le verra ultérieurement, certaines observations contredisent cette idée. Ainsi on constate que la chute d'un corps ne s'effectue pas rigoureusement le long de la verticale : on observe une déviation vers l'est. De même, l'analyse du mouvement du pendule de Foucault conduit à un phénomène d'oscillations dans un plan. Or la réalisation de l'expérience montre que le plan d'oscillations tourne. Ces deux exemples illustrent l'oubli majeur qui est fait en considérant que le référentiel terrestre est galiléen : on néglige la rotation de la Terre sur elle-même. Par conséquent, le référentiel terrestre n'est pas galiléen.

¹ Elle sera remplacée par la transformation de Lorentz en mécanique relativiste.

² Le détail de l'analyse a été faite en première période en application du principe fondamental de la dynamique.

Ce qui a conduit à rejeter le référentiel terrestre comme référentiel galiléen est la non prise en compte de la rotation de la Terre. L'idée naturelle est alors de considérer le *référentiel géocentrique* dont l'origine est le centre de la Terre et dont les axes ont des directions fixes (qui sont celles du référentiel de Copernic défini ultérieurement). Dans ce référentiel, les problèmes liés à la rotation de la Terre disparaissent et on obtient une description correcte de la déviation vers l'est lors de la chute libre d'un point matériel et du mouvement du pendule de Foucault par exemple.

Mais, lorsqu'on étudie le phénomène des marées, on constate un écart entre la prédiction obtenue à partir des lois de Newton et les observations réelles. En effet, en utilisant le référentiel géocentrique, on a oublié de tenir compte du fait que la Terre tourne autour du Soleil. L'origine du référentiel géocentrique n'a pas un mouvement de translation uniforme, mais présente un mouvement accéléré par rapport au *référentiel de Copernic* dont l'origine est confondue avec le centre de masse du système solaire et dont les axes dirigés vers trois étoiles fixes très éloignées.

Un référentiel aux propriétés très proches du référentiel de Copernic est le *référentiel de Kepler* dont l'origine est le centre du Soleil au lieu d'être le centre de masse du système solaire.

Les référentiels de Copernic ou de Kepler sont-ils galiléens ? En toute rigueur, la réponse est négative. En effet, le système solaire appartient à une galaxie, la Voie Lactée, et est soumis à l'attraction gravitationnelle des autres étoiles de la Galaxie. Le Soleil subit donc une accélération par rapport à un référentiel dont le centre serait le centre de gravité de la galaxie. On pourrait alors se dire que ce référentiel est galiléen, mais ce serait oublier que la Voie Lactée n'est qu'une galaxie au milieu d'un amas de galaxies et qu'elle subit par conséquent un mouvement accéléré par rapport à un référentiel centré sur l'ensemble des galaxies. On peut continuer ainsi le raisonnement longtemps sans pouvoir proposer un référentiel qui serait LE référentiel galiléen.

Il s'avère donc impossible de mettre en évidence un référentiel rigoureusement galiléen. On pourrait alors se dire que les lois de Newton sont des principes non recevables et chercher d'autres principes pour expliquer les phénomènes mécaniques. En fait, il n'en est rien et cela donne toute sa force au principe d'inertie. En effet, pour pouvoir considérer qu'un référentiel est en première approximation galiléen, il suffit de vérifier que les lois de Newton permettent d'interpréter correctement le comportement des points matériels étudiés. Ainsi, pour reprendre les exemples précédents, on pourra considérer le référentiel terrestre comme galiléen pour la chute d'un objet tant que la déviation vers l'est pourra être négligée. De même, la déviation vers l'est ou le comportement du pendule de Foucault pourront être traités dans le référentiel géocentrique considéré comme galiléen avec une bonne approximation. Quant au phénomène des marées, l'approximation sera le plus souvent satisfaisante en considérant que le référentiel de Copernic est galiléen.

En conclusion, on considérera un référentiel comme galiléen tant qu'on pourra appliquer les lois de Newton et en particulier le principe fondamental de la dynamique sans que les observations expérimentales n'infirment la théorie. Le référentiel de Copernic est, dans ces conditions, la meilleure approximation de référentiel galiléen qu'on puisse considérer.

1.3 Principe de relativité de Galilée

Le principe de relativité galiléenne postule que les lois de la dynamique sont les mêmes quel que soit le référentiel galiléen considéré. Pour le justifier, on va montrer que le principe fondamental de la dynamique se formule de la même manière dans tout référentiel galiléen.

Soient deux référentiels galiléens $\mathcal{R}$ et $\mathcal{R}'$.

Le principe fondamental de la dynamique ne faisant intervenir que l'accélération comme grandeur cinématique (dans le cas d'un système à masse constante), il suffit d'obtenir l'égalité des accélérations pour avoir la même formulation dans deux référentiels.

Or la loi de composition des accélérations pour passer d'un référentiel $\mathcal{R}$ à un référentiel $\mathcal{R}'$ s'écrit :

$$\vec{a}_{/\mathcal{R}}(M) = \vec{a}_{/\mathcal{R}'}(M) + \vec{a}_e(M) + \vec{a}_c(M)$$

avec $\vec{a}_e(M)$ l'accélération d'entraînement ou accélération du point coïncidant et $\vec{a}_c(M) = 2\vec{\Omega} \wedge \vec{v}_{/\mathcal{R}'}(M)$ l'accélération complémentaire ou de Coriolis. Or $\mathcal{R}$ et $\mathcal{R}'$ sont galiléens : ils sont en translation rectiligne uniforme l'un par rapport à l'autre. Le mouvement du point coïncidant est un mouvement rectiligne et uniforme donc $\vec{a}_e(M) = \vec{0}$. D'autre part, $\vec{\Omega} = \vec{0}$ donc $\vec{a}_c(M) = \vec{0}$. Les accélérations d'entraînement et de Coriolis sont nulles, d'où la relation :

$$\vec{a}_{/\mathcal{R}}(M) = \vec{a}_{/\mathcal{R}'}(M)$$

Ceci permet d'écrire le principe fondamental de la dynamique dans les deux référentiels galiléens :

$$m\vec{a}_{/\mathcal{R}}(M) = m\vec{a}_{/\mathcal{R}'}(M) = \sum_i \vec{f}_i$$

Sur le principe fondamental de la dynamique, on a vérifié ainsi le principe de relativité de Galilée selon lequel les lois de la mécanique sont invariantes par changement de référentiel galiléen, à savoir lors du passage d'un référentiel galiléen à un référentiel en translation uniforme par rapport à celui-ci.

On peut noter que ce principe dépasse largement le cadre de la mécanique : il concerne tout phénomène physique qui est donc décrit de la même manière dans

tout référentiel galiléen. C'est justement la remise en cause de ce principe par l'électromagnétisme qui a conduit à l'élaboration de la mécanique relativiste : l'existence d'une vitesse maximale égale à la vitesse de la lumière mise en évidence par la propagation des ondes électromagnétiques met en défaut la loi de composition des vitesses classiques qui a été vue au chapitre précédent.

2. Expression du principe fondamental de la dynamique en référentiel non galiléen

2.1 Position du problème

Les trois lois de Newton (principe d'inertie, principe fondamental de la dynamique et principe des actions réciproques) sont valables dans un référentiel galiléen dont le principe d'inertie postule l'existence et donne une définition.

Cependant, les référentiels galiléens ne sont pas les seuls référentiels utilisables et il est parfois intéressant de travailler dans un référentiel non galiléen. Par exemple, lors du freinage d'un véhicule, il est utile d'analyser le mouvement relatif des divers objets situés à l'intérieur du véhicule afin de pouvoir en déterminer la meilleure configuration intérieure. Le véhicule n'est pas animé d'un mouvement rectiligne uniforme par rapport au référentiel terrestre qu'on considère ici comme galiléen. Pour l'instant, les principes énoncés ne permettent pas de résoudre le problème directement dans le référentiel lié au véhicule.

On peut noter que le principe des actions réciproques ne dépend pas du référentiel dans lequel on se place et ne pose aucune difficulté. On va donc formuler le principe fondamental de la dynamique en référentiel non galiléen.

2.2 Principe fondamental de la dynamique et forces d'inertie

Soient $\mathcal{R}$ un référentiel galiléen et $\mathcal{R}'$ un référentiel non galiléen.

Lorsqu'on a étudié le problème du changement de référentiel, les deux référentiels envisagés étaient *a priori* interchangeables. Ici ce n'est plus le cas : le seul dans lequel on sait écrire le principe fondamental de la dynamique est $\mathcal{R}$ le référentiel galiléen. Ce référentiel devient de ce fait le référentiel absolu. On a donc :

$$m \vec{a}_{/\mathcal{R}}(M) = \sum_j \vec{f}_j \quad (27.1)$$

La loi de composition des accélérations s'écrit :

$$\vec{a}_{/\mathcal{R}}(M) = \vec{a}_{/\mathcal{R}'}(M) + \vec{a}_e(M) + \vec{a}_c(M)$$

La relation (27.1) devient :

$$m \vec{a}_{/\mathcal{R}'}(M) + m \vec{a}_e(M) + m \vec{a}_c(M) = \sum_j \vec{f}_j$$

ce qui permet d'exprimer le principe fondamental de la dynamique dans un référentiel non galiléen sous la forme :

$$m \vec{a}_{/\mathcal{R}'}(M) = \sum_j \vec{f}_j - m \vec{a}_e(M) - m \vec{a}_c(M)$$

On peut alors interpréter les termes $-m \vec{a}_e(M)$ et $-m \vec{a}_c(M)$ comme des forces dites *forces d'inertie*. Le terme $-m \vec{a}_e(M)$ correspond à la force d'inertie d'entraînement notée $\vec{f}_{ie}$ tandis que $-m \vec{a}_c(M)$ à la force d'inertie complémentaire ou de Coriolis notée $\vec{f}_{ic}$. On notera la présence du signe $-$ dans l'expression des forces d'inertie.

En notant $\vec{f}_{ie} = -m \vec{a}_e$ et $\vec{f}_{ic} = -m \vec{a}_c$ les forces d'inertie, le principe fondamental de la dynamique s'écrit :

$$m \vec{a}_{/\mathcal{R}'}(M) = \sum_j \vec{f}_j + \vec{f}_{ie} + \vec{f}_{ic}$$

Les forces d'inertie ne sont pas dues à des interactions : en ce sens, ce ne sont pas des forces au même titre que celles déjà présentées. Cependant elles modifient le mouvement par rapport à un référentiel non galiléen ; c'est pourquoi on peut les considérer comme des forces.

On peut noter que les accélérations d'entraînement et de Coriolis sont nulles si $\mathcal{R}'$ est galiléen, ce qui permet de retrouver l'universalité du principe fondamental de la dynamique pour tout référentiel galiléen.

2.3 Cas où $\mathcal{R}'$ est en translation par rapport à $\mathcal{R}$

On a établi dans ce cas que :

$$\begin{cases} \vec{a}_c(M) = \vec{0} \\ \vec{a}_e(M) = \frac{d^2 \overrightarrow{OO'}}{dt^2} \end{cases}$$

On en déduit que la force d'inertie de Coriolis est nulle et que la force d'inertie d'entraînement vaut :

$$\vec{f}_{ie} = -m \frac{d^2 \overrightarrow{OO'}}{dt^2}$$

2.4 Cas où $\mathcal{R}'$ est en rotation par rapport à $\mathcal{R}$

De la même façon, on a pour une rotation pure :

$$\begin{cases} \vec{a}_c(M) = 2 \vec{\Omega} \wedge \vec{v}_{/\mathcal{R}'}(M) \\ \vec{a}_e(M) = \frac{d\vec{\Omega}}{dt} \wedge \overrightarrow{O'M} - \Omega^2 \overrightarrow{HM} \end{cases}$$

en notant H la projection de M sur l'axe de rotation. D'où on déduit les expressions des forces d'inertie :

$$\begin{cases} \vec{f}_{ic} = -2m\vec{\Omega} \wedge \vec{v}_{/\mathcal{R}'}(M) \\ \vec{f}_{ie} = -m\frac{d\vec{\Omega}}{dt} \wedge \vec{O'M} + m\Omega^2\vec{HM} \end{cases}$$

Dans le cas particulier où la rotation est uniforme, la force d'inertie d'entraînement se réduit à la force dite centrifuge :

$$\vec{f}_{ie} = m\Omega^2\vec{HM}$$

qu'on retiendra sous cette forme.

2.5 Solutions pratiques

En pratique, deux solutions sont possibles pour résoudre un problème de mécanique où on cherche un mouvement relatif :

- on traite le problème dans le référentiel galiléen en appliquant le principe fondamental de la dynamique puis on effectue le changement de référentiel pour obtenir le mouvement relatif,
- on estime les forces d'inertie à partir des relations précédentes dans le référentiel mobile non galiléen et on applique alors le principe fondamental de la dynamique dans le référentiel non galiléen par :

$$m\vec{a}_{/\mathcal{R}'}(M) = \sum_i \vec{f}_i + \vec{f}_{ie} + \vec{f}_{ic}$$

Les deux approches sont rigoureusement équivalentes ; seule change la manière d'envisager la résolution.

3. Théorème du moment cinétique en référentiel non galiléen

En référentiel non galiléen $\mathcal{R}'$, le théorème s'écrit de la même façon en tenant compte des forces d'inertie d'entraînement et de Coriolis.

Ce qui se traduit mathématiquement :

$$\frac{d\vec{L}_{O/\mathcal{R}'}(M)}{dt} = \sum_j \vec{\mathcal{M}}_O(\vec{f}_j) + \vec{\mathcal{M}}_O(\vec{f}_{ie}) + \vec{\mathcal{M}}_O(\vec{f}_{ic})$$

où O est un point fixe de $\mathcal{R}'$.

En effet, en reprenant la démonstration faite dans le cas d'un référentiel galiléen, on obtient :

$$\begin{aligned}
 \frac{d\vec{L}_{O/\mathcal{R}'}(M)}{dt} &= \vec{OM} \wedge \frac{d(m\vec{v}_{/\mathcal{R}'}(M))}{dt} \\
 &= \vec{OM} \wedge \left(\sum_j \vec{f}_j + \vec{f}_{ie} + \vec{f}_{ic} \right) \\
 &= \sum_j \vec{OM} \wedge \vec{f}_j + \vec{OM} \wedge \vec{f}_{ie} + \vec{OM} \wedge \vec{f}_{ic} \\
 &= \sum_j \vec{\mathcal{M}}_O(\vec{f}_j) + \vec{\mathcal{M}}_O(\vec{f}_{ie}) + \vec{\mathcal{M}}_O(\vec{f}_{ic})
 \end{aligned}$$

4. Théorème de l'énergie cinétique en référentiel non galiléen

Dans un référentiel non galiléen, le théorème de l'énergie cinétique s'énonce par :

La variation d'énergie cinétique d'un point matériel entre deux instants est égale au travail des forces qui s'exercent sur ce point entre les deux instants considérés, y compris les forces d'inertie.

En effet, la seule chose qui change dans la démonstration de ce théorème est l'énoncé du principe fondamental de la dynamique :

$$m \frac{d\vec{v}_{/\mathcal{R}'}(M)}{dt} = \sum_j \vec{f}_j + \vec{f}_{ie} + \vec{f}_{ic}$$

Le calcul du travail élémentaire dans le référentiel $\mathcal{R}'$ donne alors :

$$\begin{aligned}
 \delta W_{/\mathcal{R}'} &= \left(\sum_j \vec{f}_j + \vec{f}_{ie} + \vec{f}_{ic} \right) \cdot \vec{v}_{/\mathcal{R}'}(M) dt = \delta W'_{/\mathcal{R}'} + \delta W_{ie} + \delta W_{ic} \\
 &= m \frac{d\vec{v}_{/\mathcal{R}'}(M)}{dt} \cdot \vec{v}_{/\mathcal{R}'}(M) dt = d \left(\frac{1}{2} m v_{/\mathcal{R}'}^2(M) \right) \\
 &= dE_{c/\mathcal{R}'}
 \end{aligned}$$

On peut noter cependant que le travail des forces de Coriolis est nul ; en effet ces forces s'écrivent :

$$\vec{f}_{ic} = -2m\vec{\Omega} \wedge \vec{v}_{/\mathcal{R}'}(M)$$

$\vec{f}_{ic}$ est donc perpendiculaire au vecteur vitesse et sa puissance est nulle. On peut retenir le résultat général suivant : la force d'inertie de Coriolis ne travaille pas. Dans

le théorème de l'énergie cinétique en référentiel non galiléen, il suffira donc de tenir compte, en plus du travail des forces noté $\delta W'_{/\mathcal{R}'}$, de celui de la force d'inertie d'entraînement :

$$dEc_{/\mathcal{R}'} = \delta W'_{/\mathcal{R}'} + \delta W_{ie}$$

5. Exemple de résolution d'un problème dans un référentiel non galiléen

On s'intéresse au mouvement d'un point matériel M de masse m astreint à glisser sans frottement le long d'une barre faisant un angle α avec la verticale et tournant autour de la verticale à vitesse angulaire Ω constante.

1. Système : masse m .
2. Référentiel $\mathcal{R}'$ non galiléen lié à la barre en rotation.
3. Bilan des forces :
 - poids $m\vec{g}$,
 - réaction de la barre $\vec{R}$ perpendiculaire à la barre du fait de l'absence de frottement (on ne la représente pas sur la figure ci-contre : elle a *a priori* une composante sur $\vec{u}'$ et sur $\vec{u} \wedge \vec{u}'$),
 - force d'inertie de Coriolis :

$$\vec{f}_{ic} = -2m\vec{\Omega} \wedge \vec{v}_{/\mathcal{R}'}(M)$$

- force d'inertie d'entraînement :

$$\vec{f}_{ie} = m\Omega^2 \overrightarrow{HM}$$

car $\mathcal{R}'$ est en rotation uniforme par rapport au référentiel du laboratoire galiléen.

4. Résolution : on va appliquer le principe fondamental de la dynamique dans le référentiel $\mathcal{R}'$ non galiléen :

$$m\vec{a}_{/\mathcal{R}'}(M) = m\vec{g} + \vec{R} - 2m\vec{\Omega} \wedge \vec{v}_{/\mathcal{R}'}(M) + m\Omega^2 \overrightarrow{HM}$$

On repère la position de M par r sa distance à O (le long de la barre).

On obtient l'équation du mouvement en projetant le principe fondamental de la dynamique suivant $\vec{u}$ pour s'affranchir de la réaction de la barre qui est inconnue :

$$m\ddot{r} = -mg \cos \alpha + m\Omega^2 HM \sin \alpha$$

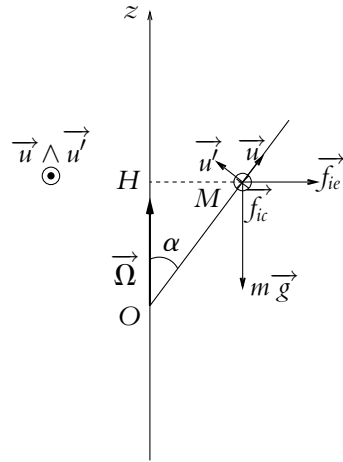


Figure 27.1 Point matériel sur une barre en rotation.

soit, compte tenu du fait que $HM = r \sin \alpha$:

$$\ddot{r} - (\Omega^2 \sin^2 \alpha) r = -g \cos \alpha$$

On peut remarquer que la force d'inertie de Coriolis n'intervient pas non plus : elle est dirigée selon $\vec{u} \wedge \vec{u}'$.

La solution de cette équation différentielle est donc la somme d'une solution particulière constante : $r_1 = \frac{g \cos \alpha}{\Omega^2 \sin^2 \alpha}$ et de la solution générale de l'équation homogène associée : $r_2 = A \exp(-\sin \alpha t) + B \exp(\Omega \sin \alpha t)$. Si on choisit comme condition initiale $r = \frac{g \cos \alpha}{\Omega^2 \sin^2 \alpha} + r_0$ et $\dot{r} = 0$ on a :

$$r = \frac{g \cos \alpha}{\Omega^2 \sin^2 \alpha} + r_0 \cosh(\Omega \sin \alpha t)$$

Si $r_0 > 0$, la masse M quitte la barre puisque r tend vers l'infini quand le temps tend également vers l'infini.

Si $r_0 < 0$, la masse M se dirige vers le point O et est finalement stoppée par l'axe.

Ayant déterminé le mouvement de M , on peut déterminer maintenant la réaction de la barre en projetant le principe fondamental de la dynamique sur $\vec{u}'$:

$$0 = R_1 - mg \sin \alpha - m\Omega^2 r \sin \alpha \cos \alpha$$

et sur $\vec{u} \wedge \vec{u}'$:

$$0 = R_2 - 2m\Omega \sin \alpha \dot{r}$$

soit

$$\begin{cases} R_1 = \frac{mg}{\sin \alpha} + m\Omega^2 \sin \alpha \cos \alpha r_0 \cosh(\Omega \sin \alpha t) \\ R_2 = 2m\Omega^2 \sin^2 \alpha r_0 \sinh(\Omega \sin \alpha t) \end{cases}$$

avec $\vec{R} = R_1 \vec{u}' + R_2 (\vec{u} \wedge \vec{u}')$.

On remarque que la réaction a deux composantes : elle est perpendiculaire à une droite donc contenue dans le plan perpendiculaire à cette droite.

A. Application directe du cours

1. Singe et pèse-personne

Un singe ayant trouvé un pèse-personne le teste dans un balançoire se déplaçant verticalement : il se place dessus et regarde l'indication de l'appareil. Lors du mouvement vers le haut, on distingue trois phases :

- une phase d'accélération constante : $a_z = 3 \text{ m.s}^{-2}$,
- une phase à vitesse v_z constante,
- une phase de décélération constante : $a_z = -3 \text{ m.s}^{-2}$.

1. Déterminer la réaction du pèse-personne sur le singe en fonction de m , g et a_z .
2. En déduire les valeurs indiquées par le pèse-personne au cours des différentes phases.

2. La fusée du Professeur Tournesol

Dans l'album de Tintin, "On a marché sur la Lune", l'accélération de la fusée maintient une pesanteur égale à la pesanteur terrestre d'après le Professeur Tournesol. On considère le référentiel géocentrique $\mathcal{R}_G$ galiléen. On rappelle la force de Newton exercée par un astre A de masse m_A sur un corps M de masse m : $\vec{F} = -Gm_A m \frac{\vec{AM}}{AM^3}$.

1. Lors de la première phase du voyage, la fusée s'éloigne de la Terre. Quelle accélération $\vec{a}_e$ la fusée doit-elle avoir par rapport au référentiel $\mathcal{R}_G$ pour que les passagers ressentent la même pesanteur qu'à la surface de la Terre ? On note R_T le rayon de la Terre et l'on considère que l'attraction lunaire est négligeable.
2. Dans la deuxième phase du voyage, la fusée s'est retournée et n'est plus soumise qu'à l'attraction lunaire. Quelle doit-être la nouvelle accélération pour que les passagers aient l'impression d'être sur Terre ? On considèrera que le référentiel lié à la Lune est aussi galiléen.

3. Jeu de billes

Pendant le trajet de départ en vacances en camping-car, un enfant observe une bille qu'il a laissé sur la table à l'arrière.

1. Tout à coup, il voit la bille partir vers l'arrière. Comment expliquer ce mouvement ?
2. À un autre moment, le conducteur accélère avec une accélération constante $\vec{a}$, quel est alors le mouvement de la bille dans le camping-car si elle est initialement au repos ?
3. Enfin, le conducteur prend un virage de rayon R assez grand devant la taille du véhicule, à vitesse constante V . Quel est alors le mouvement de la bille dans le camping-car si elle est initialement au repos ?

4. Expérience du manège inertiel

Si vous allez à la Cité des Sciences à la Villette, vous pourrez monter dans un manège inertiel. Lorsque l'on est entré, les portes se referment et le manège se met en mouvement par rapport au référentiel terrestre : mouvement de rotation autour d'un axe vertical, à la vitesse angulaire constante ω , dans le sens trigonométrique pour un observateur extérieur. L'une des expériences

consiste à se placer au centre du manège et à essayer d'atteindre une cible située en face de soi, à une distance $d = 3$ m du centre, avec une balle en mousse.

1. Définir le référentiel absolu et galiléen $\mathcal{R}$ et le référentiel relatif $\mathcal{R}'$ nécessaires à l'étude. Quel est le mouvement de $\mathcal{R}'$ par rapport à $\mathcal{R}$?

Les questions (2) et (3) ne nécessitent pas de calcul :

2. Lors du premier lancé, le tireur vise la cible. Est-ce que la balle atteint la cible ? Est-ce que son point d'impact est à gauche ou à droite de la cible ? (justifier la réponse).

3. Lors du deuxième tir, le tireur n'ayant pas atteint la cible la première fois, il pense qu'il suffit de tirer plus fort. A-t-il raison ? A-t-il tort ? Ne peut-il conclure ?

La distance entre les points de départ et d'arrivée étant faible, on suppose que la vitesse initiale v_0 est suffisante pour que le mouvement ait lieu dans un plan horizontal. Soit (O, x, y, z) un repère lié au manège avec Oz vertical, la cible est sur l'axe Ox , et le tireur au centre.

4. L'hypothèse précédente conduit à négliger l'effet d'une force, laquelle ?

5. Établir les deux équations différentielles du mouvement dans le référentiel lié au manège.

6. En déduire les équations paramétriques de la trajectoire $x(t)$ et $y(t)$, en fonction de la vitesse initiale v_0 suivant Ox .

7. Déterminer le point d'impact sur le mur et commentez vos réponses aux questions (2) et (3).

5. Adhérence d'une masse sur un plateau oscillant

Un point matériel de masse m est posé sur un plateau animé d'un mouvement vertical sinusoïdal $z = a \cos \omega t$.

À quelle condition la masse ne quitte-t-elle pas le plateau ?

A.N. : $a = 5$ cm, $g = 10$ m.s⁻².

Montrer l'existence d'une fréquence critique et calculer sa valeur.

6. Portière mal fermée

La portière d'une voiture est restée entrouverte au moment où le véhicule se met à freiner avec une décélération constante a . La portière est modélisée par une masse m située à une distance l de la charnière au bout d'une tige de masse négligeable.

1. Écrire l'équation différentielle du mouvement.

2. Déterminer la vitesse de la porte pour un angle de 90° .

3. On étudie maintenant la phase d'accélération. Déterminer la relation entre θ_0 l'angle initial de la porte avec l'aile et a pour que la porte se ferme.

7. Masselotte en rotation sur une tige, d'après ICNA 2005

Une masselotte ponctuelle M , de masse m , peut glisser sans frottement sur une tige perpendiculaire en O à un axe vertical Oz . L'axe z est entraîné par un moteur qui fait tourner la tige

(T) à la vitesse angulaire constante ω dans le plan horizontal x_0Oy_0 d'un repère galiléen fixe orthonormé direct (O, x_0, y_0, z) . Soient $(\vec{u}_{x_0}, \vec{u}_{y_0}, \vec{u}_z)$ les vecteurs unitaires de chacun des axes du repère galiléen $\mathcal{R}_0$, et $(\vec{u}_x, \vec{u}_y, \vec{u}_z)$ les vecteurs unitaires du repère orthonormé direct $\mathcal{R}(O, x, y, z)$ lié à (T).

La tige (T) est portée par l'axe Ox . La masselotte est repérée par ses coordonnées cylindropolaires (ρ, φ, z) liées à $\mathcal{R}_0$.

À l'instant initial $t = 0$, la tige (T) est confondue avec l'axe Ox_0 , et la masselotte est lancée depuis le point O avec une vitesse initiale $\vec{v}_0 = v_0 \vec{u}_x$ où $v_0 > 0$. On note g l'intensité du champ de pesanteur terrestre.

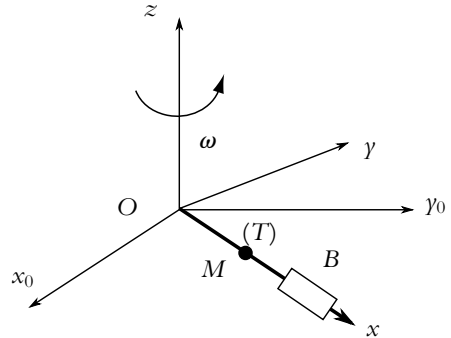


Figure 27.2

1. a) Effectuer le bilan des forces dans le repère $\mathcal{R}$ et exprimer chacune des forces en fonction des données.
- b) Exprimer la relation fondamentale de la dynamique dans $\mathcal{R}$ en projection sur $(\vec{u}_x, \vec{u}_y, \vec{u}_z)$ en fonction des données.
2. Déterminer l'équation horaire du mouvement $\rho(t)$.
3. Par composition des vitesses, déterminer la vitesse de M dans $\mathcal{R}_0$.
4. À la distance D de O , on place au point B de la tige une butée (B) solidaire de (T). À l'instant t_B , la masselotte vient buter sur (B). Si la tige effectue un tour complet en 16 s, avec $v_0 = 0,393 \text{ m.s}^{-1}$ et $D = 2,3 \text{ m}$, calculer t_B .
5. Quelle est la trajectoire décrite par M dans $\mathcal{R}_0$ avant et après t_B ?
6. Lorsque la masselotte est en butée sur (B), exprimer la réaction $\vec{R}_H$ de (B) sur M .

8. Mouvement d'une masse dans un cerceau horizontal

Un cerceau tourne autour d'un axe vertical perpendiculaire à son plan à vitesse angulaire ω constante comme indiqué sur la figure. Une masse m glisse sans frottement à l'intérieur du cerceau. Déterminer l'équation du mouvement dans le cas où il n'y a pas de frottements.

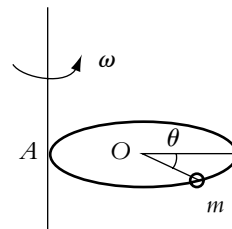


Figure 27.3

B. Exercices et problèmes

1. Masse astreinte sur un axe et reliée à un autre par un ressort

La masse m se déplace le long de l'axe Ox' sans frottement. L'axe Ox' tourne autour de Oz à vitesse angulaire ω constante.

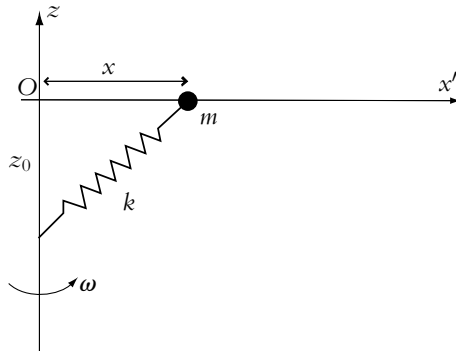


Figure 27.4

1. Exprimer les accélérations de Coriolis et d'entraînement de la barre Ox' .
2. Etablir les équations du mouvement. On pourra soit projeter le principe fondamental de la dynamique soit appliquer un raisonnement énergétique. On justifiera la validité de la méthode utilisée.
3. Déterminer les positions d'équilibre de m en discutant suivant les valeurs de ω .
4. Discuter la stabilité de la position d'équilibre $x = 0$ en fonction des valeurs de ω .
5. Même question pour l'éventuelle autre position d'équilibre.

2. Pendule en translation, d'après ENAC 2007

On désigne par $\mathcal{R}'(O', x', y', z')$ un référentiel d'origine O' dont les axes orthogonaux $O'x'$, $O'y'$ et $O'z'$ sont respectivement parallèles aux axes Ox , Oy et Oz d'un référentiel $\mathcal{R}(O, x, y, z)$ que l'on supposera galiléen. Un pendule simple est constitué d'un point matériel P de masse m , suspendu à l'origine O' de $\mathcal{R}'$ par un fil sans masse ni raideur et de longueur ℓ . On note θ l'angle que fait le fil, que l'on supposera constamment tendu, avec la verticale Oy de $\mathcal{R}$ (figure 27.5).

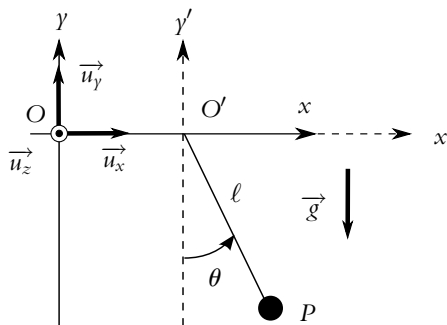


Figure 27.5

Dans un premier temps l'origine O' de $\mathcal{R}'$ reste fixe et confondue avec l'origine O de $\mathcal{R}$.

1. Quelle doit être la longueur ℓ du fil pour que la période des petits mouvements soit $T_0 = 1$ s. On prendra pour norme de l'accélération de la pesanteur $\vec{g} = -g\vec{u}_y$, la valeur $g = 9,8 \text{ m.s}^{-2}$?

- A) $\ell = 1,141 \text{ m}$ B) $\ell = 1,714 \text{ m}$
C) $\ell = 1,312 \text{ m}$ D) $\ell = 0,248 \text{ m}$

2. Le référentiel $\mathcal{R}'$ est maintenant animé d'un mouvement de translation rectiligne uniformément accéléré, d'accélération constante $\vec{a} = a\vec{u}_x$.

Calculer le moment $\vec{\mathcal{M}}'_O(\vec{F}_{ie})$ par rapport au point O' de la force d'inertie d'entraînement $\vec{F}_{ie}$ qui s'applique au point P dans le référentiel $\mathcal{R}'$.

- A) $\vec{\mathcal{M}}'_O(\vec{F}_{ie}) = -m\ell a \cos \theta \vec{u}_z$ B) $\vec{\mathcal{M}}'_O(\vec{F}_{ie}) = m\ell a (\cos \theta - \sin \theta) \vec{u}_z$
C) $\vec{\mathcal{M}}'_O(\vec{F}_{ie}) = m\ell a (\cos \theta + \sin \theta) \vec{u}_x$ D) $\vec{\mathcal{M}}'_O(\vec{F}_{ie}) = -m\ell a \sin \theta \vec{u}_y$

3. Calculer le moment $\vec{\mathcal{M}}'_O(\vec{F}_{ic})$ par rapport au point O' de la force d'inertie de Coriolis $\vec{F}_{ic}$ qui s'applique au point P dans le référentiel $\mathcal{R}'$.

- A) $\vec{\mathcal{M}}'_O(\vec{F}_{ic}) = -m\ell^2 a \vec{u}_z$ B) $\vec{\mathcal{M}}'_O(\vec{F}_{ic}) = m\ell^2 \ddot{\theta} \vec{u}_x$
C) $\vec{\mathcal{M}}'_O(\vec{F}_{ic}) = -m\ell \dot{\theta} \cos \theta \vec{u}_z$ D) $\vec{\mathcal{M}}'_O(\vec{F}_{ic}) = \vec{0}$

4. Dédurre du théorème du moment cinétique appliqué en O' dans $\mathcal{R}'$ au point matériel P l'équation différentielle à laquelle obéit l'angle θ .

- A) $\ddot{\theta} = \frac{g}{\ell} \cos \theta + \frac{a}{\ell} \sin \theta$ B) $\ddot{\theta} = -\frac{a}{\ell} \cos \theta + \frac{g}{\ell} \sin \theta$
C) $\ddot{\theta} = -\frac{a}{\ell} \sin \theta + \frac{g}{\ell} \cos \theta$ D) $\ddot{\theta} = -\frac{g}{\ell} \sin \theta - \frac{a}{\ell} \cos \theta$

5. Déterminer la valeur de l'angle θ_0 correspondant à la position d'équilibre du pendule.

- A) $\theta_0 = -\arctan \frac{a}{g}$ B) $\theta_0 = \arctan \frac{a}{g}$
C) $\theta_0 = \arctan \frac{g}{a}$ D) $\theta_0 = -\arctan \frac{g}{a}$

6. Exprimer la période T des petits mouvements autour de la position d'équilibre θ_0 en fonction de ℓ , a et g .

- A) $T = 2\pi \sqrt{\frac{\ell a}{a^2 + g^2}}$ B) $T = 2\pi \sqrt{\frac{\ell}{\sqrt{a^2 + g^2}}}$
C) $T = 2\pi \sqrt{\frac{\ell g}{a^2 + g^2}}$ D) $T = 2\pi \sqrt{\frac{\ell}{a + g}}$

3. Étude d'un ressort dans un référentiel tournant, d'après ENSTIM 2002

Un palet M de masse m peut se mouvoir sans frottement dans le plan (O, x, y) horizontal (table à coussin d'air par exemple). Il est accroché à l'extrémité d'un ressort de longueur à vide ℓ_0 , de raideur k , dont l'autre extrémité est fixée en O . Le champ de pesanteur est suivant la verticale Oz : $\vec{g} = g\vec{u}_z$.

On définit deux référentiels :

- Le référentiel $\mathcal{R}$ galiléen associé au repère fixe $(0, \vec{u}_x, \vec{u}_y, \vec{u}_z)$.
- Le référentiel $\mathcal{R}'$ associé au repère $(0, \vec{u}_r, \vec{u}_\theta, \vec{u}_z)$, en rotation par rapport à $\mathcal{R}$ avec un vecteur rotation $\vec{\omega} = \omega\vec{u}_z$.

On lance la particule d'un point M_0 , tel que $\vec{OM}_0 = \ell_1\vec{u}_x$, avec une vitesse initiale $\vec{v}_0 = \ell_1\omega\vec{u}_y$.

Le mouvement est étudié dans le référentiel $\mathcal{R}'$. La position de M est repérée dans la base $(\vec{u}_r, \vec{u}_\theta)$ par $\vec{OM} = r\vec{u}_r$ sachant qu'à $t = 0$, $\vec{u}_r = \vec{u}_x$.

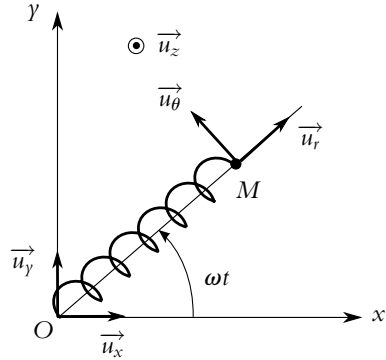


Figure 27.6

1. Préciser les expressions vectorielles des forces d'inertie dans la base $(\vec{u}_r, \vec{u}_\theta, \vec{u}_z)$.
2. Montrer que la force d'inertie d'entraînement dérive d'une énergie potentielle que l'on précisera.
3. En est-il de même pour la force d'inertie de Coriolis ou complémentaire ?
4. Déterminer l'énergie potentielle totale E_p . Tracer l'allure de $E_p(r)$. On distinguera les 3 cas possibles selon la valeur de ω .
5. Déterminer la longueur ℓ_2 correspondant à la position d'équilibre dans le référentiel $\mathcal{R}'$.
6. À quelle condition sur ω le résultat est-il possible ? Cet équilibre est-il stable ?
7. Quel est alors le mouvement dans le référentiel $\mathcal{R}$?
8. Comparer ℓ_2 à ℓ_1 .

4. Équilibre dans un bol en rotation

Une bille assimilée à un point matériel M , de masse m est astreinte à se déplacer, sans frottement, sur la surface intérieure d'une demi-sphère creuse (S) . Cette surface tourne uniformément, à la vitesse angulaire ω constante, autour de son axe de révolution vertical (Oz) . On appelle r_0 la rayon intérieur de la sphère, et ω_0 la grandeur $\sqrt{2g/r_0}$, g étant l'accélération de la pesanteur. On note $\mathcal{R}_\mathcal{L}$ le référentiel du laboratoire galiléen et $\mathcal{R}_\mathcal{S}$ le référentiel lié à (S) .

1. Déterminer l'énergie potentielle de M dans le référentiel $\mathcal{R}_\mathcal{S}$. On utilisera un repère cylindrique d'axe Oz , O centre de la sphère complète. Montrer qu'elle peut se mettre sous la forme :

$$E_p = \frac{1}{2}m\omega^2 z \left(\left(\frac{\omega_0}{\omega} \right)^2 r_0 + z \right)$$

2. Déterminer les cotes z_e des positions d'équilibre de M dans $\mathcal{R}_\mathcal{S}$ ainsi que leur stabilité en fonction de la valeur de ω_0/ω . Représenter z_e en fonction de ω/ω_0 .

3. a) Tracer la courbe $E_p(z)$ dans le cas $\omega_0/\omega = 0, 2$.

b) Un dispositif permet de lâcher la bille sans vitesse initiale dans $\mathcal{R}_S$. On note z_0 la cote initiale. Discuter qualitativement le mouvement ultérieur de la bille en fonction de son énergie mécanique.

5. Équilibre d'une bille sur une piste en rotation - Bifurcation

Un point A de masse m évolue sans frottement sur une piste circulaire de centre O et de rayon r . Ce guide tourne autour de son diamètre vertical à la vitesse angulaire constante Ω .

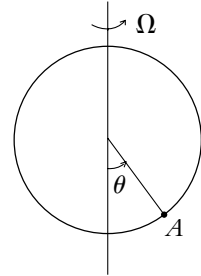


Figure 27.7

1. Quelles sont les forces auxquelles est soumis le point A dans le référentiel lié au guide en rotation par rapport à la verticale ?

2. Justifier qu'un raisonnement énergétique est adapté à ce problème.

3. Exprimer l'énergie cinétique de A dans le repère lié au guide en rotation par rapport à la verticale.

4. Donner l'expression de l'énergie potentielle de pesanteur et montrer que la force d'inertie d'entraînement dérive d'un potentiel. On prendra $\theta = \frac{\pi}{2}$ comme origine des potentiels.

5. Déterminer l'équation différentielle vérifiée par θ .

6. Quelles sont les positions d'équilibre du point A ?

7. Etudier leur stabilité.

8. Tracer les différentes positions d'équilibre en fonction de Ω . Il s'agit du diagramme de bifurcation de ce système.

6. Étude d'un sismomètre

On définit un référentiel $\mathcal{R}$ galiléen lié à l'axe Oz fixe. Une barre PX fait un angle α avec Oz (figure 27.8). Un point M de masse m , lié à un ressort (k, ℓ_0) est mobile sans frottement sur cette barre. L'accélération de la pesanteur $\vec{g}$ est uniforme et verticale suivant Oz . Le point M est lié à un système de freinage (non représenté sur la figure 27.8) qui exerce une force $\vec{F}_f = -\lambda \dot{X} \vec{u}_X$ où $\vec{u}_X$ est un vecteur unitaire selon PX . On note $X = \overline{PM}$.

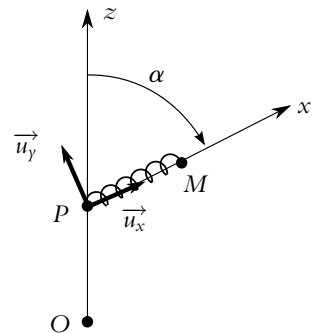


Figure 27.8

1. Déterminer la position d'équilibre X_e de M .

On considère dans la suite que le point P a un mouvement sinusoïdal dans $\mathcal{R}$: $\overline{OP} = a \cos \omega t \vec{u}_z$. On définit le référentiel $\mathcal{R}'$ lié à la barre.

2. Déterminer les forces d'inertie s'appliquant à M dans $\mathcal{R}'$. Les projeter sur $(\vec{u}_X, \vec{u}_Y)$ (figure 27.8).

3. Écrire la relation fondamentale de la dynamique dans $\mathcal{R}'$ et la projeter sur $(\vec{u}_X, \vec{u}_Y)$.

4. On pose $E = X - X_e$. Montrer que E vérifie l'équation suivante :

$$\frac{d^2 E}{dt^2} + \frac{\omega_0}{Q} \frac{dE}{dt} + \omega_0^2 E = \alpha \omega^2 \cos(\omega t)$$

où α , ω_0 et Q sont des grandeurs à déterminer en fonction des données.

5. On s'intéresse au régime forcé, c'est à dire aux solutions particulières du type $E = E_0 \cos(\omega t + \varphi)$. On utilise les notations complexes pour résoudre l'équation précédente avec $\underline{E} = E_0 \exp(i(\omega t - \varphi))$. Écrire l'équation différentielle avec les notations complexes. En déduire le module $|E|$ de $\underline{E}$.

6. On effectue le changement de variable $\Sigma = \frac{\omega_0}{\omega}$. Étudier la fonction $|E|(\Sigma)$ (dérivée et tableau de variation). En déduire les variations de $|E|$ en fonction de ω et tracer $|E|(\omega)$. Pourquoi peut-on qualifier le système de filtre passe-haut ?

7. Déduire de la courbe précédente la condition sur ω pour que le système fonctionne en sismomètre, c'est-à-dire lorsque que M reproduit les mouvements de P . Pourquoi le système de freinage est-il nécessaire ?

Cinétique d'un système de deux points matériels

28

Dans les chapitres précédents, on s'est intéressé à l'étude d'un point matériel ou d'un système modélisable par un point matériel. Désormais on ne va plus considérer un seul point matériel mais un ensemble de points matériels. Dans le cadre du programme, on se limite au cas de deux points matériels. Dans un premier temps, il est nécessaire de définir, pour un système de points, les notions qui ont été utilisées pour un point matériel à savoir : la masse, la quantité de mouvement, le moment cinétique et l'énergie cinétique. C'est l'objet de ce chapitre.

1. Masse et centre d'inertie

1.1 Système de deux points matériels

On considère un ensemble de deux points matériels A_1 et A_2 qui constitue le système étudié¹.

1.2 Masse d'un système de deux points matériels

On définit la masse du système par le scalaire

$$M = m_1 + m_2$$

On parle aussi de masse totale afin d'éviter d'éventuelles confusions. Il est à noter que cette définition signifie que la masse d'un système présente la propriété d'additivité.

¹L'ensemble des notions qui vont être définies ici peuvent facilement se généraliser au cas d'un système de N points matériels en sommant de 1 à N au lieu de 1 à 2 ou d'un système matériel continu en remplaçant les sommes discrètes par des intégrales et en définissant des volumes élémentaires autour des points. (Cf. cours de deuxième année).

1.3 Centre d'inertie

Il s'agit d'un point particulièrement important pour un système de deux points matériels. On définit le *centre d'inertie* ou encore *centre de gravité* comme le barycentre des deux points matériels affectés de leur masse. On le note habituellement G . La traduction mathématique de cette propriété est :

$$(m_1 + m_2) \overrightarrow{OG} = M \overrightarrow{OG} = m_1 \overrightarrow{OA_1} + m_2 \overrightarrow{OA_2}$$

Cette expression est valable pour tout point O et en particulier pour le point G qui peut donc également être défini par :

$$m_1 \overrightarrow{GA_1} + m_2 \overrightarrow{GA_2} = \vec{0}$$

Ces deux définitions sont équivalentes.

2. Éléments dynamiques d'un système de points matériels

2.1 Quantité de mouvement

On définit la quantité de mouvement par rapport à un référentiel $\mathcal{R}$ d'un système de deux points matériels (qu'on peut également appeler résultante cinétique) comme la somme des quantités de mouvement par rapport au référentiel $\mathcal{R}$ de chaque point matériel composant le système :

$$\vec{P}_{/\mathcal{R}} = \vec{p}_{1/\mathcal{R}} + \vec{p}_{2/\mathcal{R}} = m_1 \vec{v}_{1/\mathcal{R}} + m_2 \vec{v}_{2/\mathcal{R}}$$

où $\vec{p}_{i/\mathcal{R}} = m_i \vec{v}_{i/\mathcal{R}}$ désigne la quantité de mouvement du point matériel A_i et $\vec{v}_{i/\mathcal{R}}$ sa vitesse dans le référentiel $\mathcal{R}$.

$$\begin{aligned} \vec{P}_{/\mathcal{R}} &= m_1 \vec{v}_{1/\mathcal{R}} + m_2 \vec{v}_{2/\mathcal{R}} \\ &= m_1 \frac{d\overrightarrow{OA_1}}{dt} + m_2 \frac{d\overrightarrow{OA_2}}{dt} \quad \text{où } O \text{ est un point fixe de } R \\ &= \frac{d}{dt} \left(m_1 \overrightarrow{OA_1} + m_2 \overrightarrow{OA_2} \right) \\ &= \frac{d(M\overrightarrow{OG})}{dt} \\ &= M \vec{v}_{/\mathcal{R}}(G) \end{aligned}$$

La quantité de mouvement totale dans le référentiel $\mathcal{R}$ est égale à la quantité de mouvement d'un point fictif confondu avec le centre d'inertie où serait concentrée la

totalité de la masse du système :

$$\vec{P}_{/\mathcal{R}} = M\vec{v}_{/\mathcal{R}}(G)$$

où $\vec{v}_{/\mathcal{R}}(G)$ est la vitesse du centre d'inertie dans le référentiel $\mathcal{R}$.

2.2 Moment cinétique

a) Moment cinétique par rapport à un point O

Pour ne pas alourdir les notations, nous noterons par la suite les vitesses des points $\vec{v}_1$ et $\vec{v}_2$ sans préciser le référentiel.

Le moment cinétique d'un système de deux points matériels par rapport à un point O est défini comme la somme des moments cinétiques de chaque point matériel par rapport au même point O :

$$\vec{L}_O = \vec{L}_{O,1} + \vec{L}_{O,2} = \vec{OA}_1 \wedge m_1 \vec{v}_1 + \vec{OA}_2 \wedge m_2 \vec{v}_2$$

Il est à noter que le moment cinétique d'un système de deux points matériels par rapport à un point O dépend de ce point O. Comme on l'a fait pour un point matériel, on montre que :

$$\vec{L}_{O'} = \vec{L}_O + \vec{O'O} \wedge \vec{P} \neq \vec{L}_O$$

En effet :

$$\begin{aligned} \vec{L}_{O'} &= \vec{O'A}_1 \wedge m_1 \vec{v}_1 + \vec{O'A}_2 \wedge m_2 \vec{v}_2 \\ &= (\vec{O'O} + \vec{OA}_1) \wedge m_1 \vec{v}_1 + (\vec{O'O} + \vec{OA}_2) \wedge m_2 \vec{v}_2 \\ &= \vec{O'O} \wedge (m_1 \vec{v}_1 + m_2 \vec{v}_2) + \vec{OA}_1 \wedge m_1 \vec{v}_1 + \vec{OA}_2 \wedge m_2 \vec{v}_2 \\ &= \vec{O'O} \wedge \vec{P} + \vec{L}_O \end{aligned}$$

Il n'est pas utile de retenir cette relation : les risques d'inverser l'ordre du produit vectoriel sont grands ! En revanche, il faut être capable de l'établir rapidement.

2.3 Énergie cinétique

Elle est définie comme la somme des énergies cinétiques de chacun des deux points matériels constituant le système :

$$E_c = \frac{1}{2}m_1 v_1^2 + \frac{1}{2}m_2 v_2^2$$

en notant m_1 et v_1 respectivement la masse et la vitesse du point matériel A_1 et m_2 et v_2 celles du point matériel A_2 . L'énergie cinétique d'un système de deux points matériels, comme celle d'un point matériel, dépend du référentiel dans laquelle elle est calculée.

2.4 Éléments cinétiques d'un système de deux points matériels

Comme pour un point matériel, on a donc défini des éléments cinétiques pour le système de deux points matériels, qui sont ici :

- la quantité de mouvement ou résultante cinétique :

$$\vec{P} = m_1 \vec{v}_1 + m_2 \vec{v}_2 = M \vec{v}(G)$$

- le moment cinétique par rapport à un point O :

$$\vec{L}_O = \overrightarrow{OA_1} \wedge m_1 \vec{v}_1 + \overrightarrow{OA_2} \wedge m_2 \vec{v}_2$$

- l'énergie cinétique

$$Ec = \frac{1}{2} m_1 v_1^2 + \frac{1}{2} m_2 v_2^2$$

Ils sont tous définis par rapport à un référentiel donné $\mathcal{R}$.

3. Référentiel barycentrique et théorèmes de Koenig

On a fait apparaître un point particulier : le centre d'inertie G du système. On va ici définir un référentiel utilisant ce point particulier puis établir plusieurs résultats importants.

3.1 Référentiel barycentrique

On appelle *référentiel barycentrique* du système de deux points matériels A_1 et A_2 relativement à un référentiel galiléen $\mathcal{R}$, et on note $\mathcal{R}^*$ le référentiel animé par rapport au référentiel $\mathcal{R}$ d'un mouvement de translation à la vitesse $\vec{v}(G)$, vitesse du centre d'inertie du système.

Attention ! Le référentiel barycentrique n'est en général pas galiléen. En effet, il le serait s'il était en translation rectiligne uniforme c'est-à-dire si la vitesse $\vec{v}_G$ du centre d'inertie était constante ; ce qui n'est généralement pas le cas.

Par convention, on notera toutes les grandeurs relatives au référentiel barycentrique avec une étoile.

La définition du référentiel barycentrique signifie que ses axes ont des directions fixes au cours du temps par rapport au référentiel $\mathcal{R}$. Cela est lié à la propriété de translation.

Le cas représenté ci-dessous est le cas particulier où on prend comme directions fixes des axes celles du référentiel $\mathcal{R}$ lui-même :

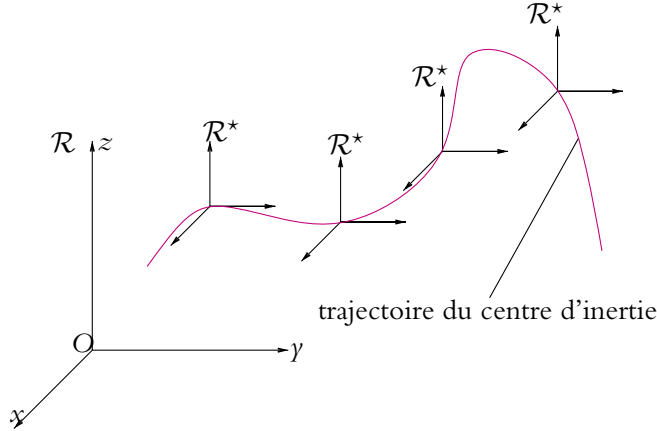


Figure 28.1 Référentiel barycentrique.

3.2 Propriétés fondamentales

a) La quantité de mouvement est nulle dans le référentiel barycentrique

On a établi que, dans tout référentiel $\mathcal{R}$:

$$m_1 \vec{v}_{1/\mathcal{R}} + m_2 \vec{v}_{2/\mathcal{R}} = M \vec{v}_{/\mathcal{R}}(G)$$

soit, en prenant comme référentiel $\mathcal{R}^*$ le référentiel barycentrique :

$$m_1 \vec{v}_1^* + m_2 \vec{v}_2^* = M \vec{v}_G^* = \vec{0}$$

puisque G est fixe dans le référentiel barycentrique.

La quantité de mouvement d'un système de deux points matériels est nulle dans le référentiel barycentrique :

$$\vec{P}^* = \vec{0}$$

b) Moment cinétique barycentrique

On a établi la relation :

$$\vec{L}_{O'} = \vec{L}_O + \vec{O'O} \wedge \vec{P}$$

dans tout référentiel $\mathcal{R}$. Donc dans le référentiel barycentrique :

$$\vec{L}_{O'}^* = \vec{L}_O^* + \vec{O'O} \wedge \vec{P}^*$$

Comme on a établi que $\vec{P}^* = \vec{0}$, on en déduit :

$$\vec{L}_{O'}^* = \vec{L}_O^*$$

Par conséquent, dans le référentiel barycentrique, le moment cinétique est indépendant du point par rapport auquel on le calcule. On parle de *moment cinétique barycentrique* et souvent (mais ce n'est pas une obligation) on choisit comme point O le centre d'inertie G . On le note $\vec{L}$.

3.3 Premier théorème de Koenig : moment cinétique dans le référentiel barycentrique

On cherche à établir la relation entre le moment cinétique barycentrique et le moment cinétique par rapport à un point O dans le référentiel $\mathcal{R}$ initial.

On part de la relation définissant le moment cinétique par rapport à un point O dans $\mathcal{R}$:

$$\vec{L}_O = \vec{OA}_1 \wedge m_1 \vec{v}_1 + \vec{OA}_2 \wedge m_2 \vec{v}_2$$

La loi de composition des vitesses pour passer du référentiel $\mathcal{R}$ au référentiel barycentrique s'écrit :

$$\vec{v}_i = \vec{v}_i^* + \vec{v}_e = \vec{v}_i^* + \vec{v}_{/\mathcal{R}}(G) \quad \text{où } i = 1 \text{ ou } 2$$

Donc

$$\begin{aligned} \vec{L}_O &= \vec{OA}_1 \wedge m_1 \vec{v}_1^* + \vec{OA}_2 \wedge m_2 \vec{v}_2^* + \vec{OA}_1 \wedge m_1 \vec{v}_{/\mathcal{R}}(G) + \vec{OA}_2 \wedge m_2 \vec{v}_{/\mathcal{R}}(G) \\ &= \vec{L}_O^* + \left(m_1 \vec{OA}_1 + m_2 \vec{OA}_2 \right) \wedge \vec{v}_{/\mathcal{R}}(G) \\ &= \vec{L}^* + M \vec{OG} \wedge \vec{v}_{/\mathcal{R}}(G) \end{aligned}$$

On démontre ainsi le premier théorème de Koenig :

Le moment cinétique d'un système de deux points matériels par rapport à un point O est égal à la somme :

- du moment cinétique par rapport à ce point O d'un point matériel fictif confondu avec le centre d'inertie où serait concentrée la totalité de la masse du système,
- du moment cinétique du système de deux points matériels par rapport au centre d'inertie exprimé dans le référentiel barycentrique.

$$\vec{L}_O = \vec{OG} \wedge M \vec{v}_{/\mathcal{R}}(G) + \vec{L}^*$$

3.4 Second théorème de Koenig : énergie cinétique dans le référentiel barycentrique

On effectue la même étude pour l'énergie cinétique :

$$E_c = \frac{1}{2}m_1v_1^2 + \frac{1}{2}m_2v_2^2$$

Or

$$\vec{v}_i = \vec{v}_{/\mathcal{R}}(G) + \vec{v}_i^*$$

$$\begin{aligned} E_c &= \frac{1}{2}m_1 \left(v_{/\mathcal{R}}^2(G) + v_1^{*2} + 2\vec{v}_{/\mathcal{R}}(G) \cdot \vec{v}_1^* \right) + \frac{1}{2}m_2 \left(v_{/\mathcal{R}}^2(G) + v_2^{*2} + 2\vec{v}_{/\mathcal{R}}(G) \cdot \vec{v}_2^* \right) \\ &= \frac{1}{2}(m_1 + m_2) v_{/\mathcal{R}}^2(G) + \frac{1}{2}m_1 v_1^{*2} + \frac{1}{2}m_2 v_2^{*2} + \vec{v}_{/\mathcal{R}}(G) \cdot (m_1 \vec{v}_1^* + m_2 \vec{v}_2^*) \\ &= \frac{1}{2}M v_{/\mathcal{R}}^2(G) + E_c^* + \vec{v}_{/\mathcal{R}}(G) \cdot \vec{0} \end{aligned}$$

On obtient donc le second théorème de Koenig :

L'énergie cinétique d'un système de deux points matériels est égale à la somme :

- de l'énergie cinétique d'un point matériel fictif confondu avec le centre d'inertie où serait concentrée la totalité de la masse du système,
- de l'énergie cinétique du système de deux points matériels par rapport au référentiel barycentrique.

$$E_c = \frac{1}{2}M v_{/\mathcal{R}}^2(G) + E_c^*$$

29

Dynamique d'un système de deux points matériels

On s'intéresse maintenant aux lois dynamiques qui permettent d'établir la nature et la trajectoire du mouvement d'un système de deux points matériels. On s'intéresse dans un premier temps aux actions engendrant le mouvement puis aux théorèmes de la résultante cinétique, du moment cinétique et de l'énergie cinétique pour un système de deux points matériels.

1. Actions extérieures et intérieures

1.1 Système de forces

On considère un ensemble de deux points matériels A_1 de masse m_1 et A_2 de masse m_2 sur lesquels s'exercent des forces. Pour chaque point matériel A_i , on désigne par $\vec{f}_i$ la somme (ou résultante) de l'ensemble des forces qui s'exercent sur ce point matériel. L'ensemble des forces $\vec{f}_i$ constitue un système de forces s'exerçant sur le système.

1.2 Forces intérieures et forces extérieures

On est toujours dans le cadre de la dynamique newtonienne, la seule différence est le nombre de points matériels à envisager : on considère désormais deux points matériels et non un seul.

Dans ce cadre de la dynamique newtonienne, la troisième loi de Newton sur les actions réciproques reste valable. Cela conduit naturellement à distinguer deux types de forces :

- les forces extérieures décrivant l'interaction entre l'un des points du système considéré ou les deux et l'extérieur,
- les forces intérieures traduisant l'interaction entre les deux points matériels du système.

On notera $\vec{f}_{i,\text{ext}}$ la résultante des forces extérieures s'exerçant sur le point matériel A_i et $\vec{f}_{ij}$ la force exercée par le point matériel A_i sur le point matériel A_j .

La troisième loi de Newton énonce que les forces $\vec{f}_{12}$ et $\vec{f}_{21}$ (en utilisant la notation qui vient d'être donnée) sont à chaque instant opposées : elles ont même droite d'action, même norme mais un sens opposé, soit

$$\vec{f}_{21} = -\vec{f}_{12}$$

Il ne faudra jamais oublier cette propriété fondamentale : les forces intérieures d'un système de deux points matériels sont opposées. Cela aura de nombreuses conséquences par la suite.

1.3 Résultante des forces

Du fait de la distinction faite au paragraphe précédent entre forces intérieures et forces extérieures, la résultante de forces $\vec{F}$ sera donnée par la somme de deux vecteurs :

- le premier est la résultante des forces extérieures,
- le second la somme des forces intérieures.

Mathématiquement cela se traduit par :

$$\vec{F} = \vec{f}_1 + \vec{f}_2 = \vec{f}_{1,\text{ext}} + \vec{f}_{2,\text{ext}} + \vec{f}_{12} + \vec{f}_{21}$$

On remarquera que $\vec{f}_i$ correspond à la résultante des forces sur le point matériel A_i que ces forces soient intérieures ou extérieures.

La somme des forces intérieures est :

$$\vec{F}_{\text{int}} = \vec{f}_{12} + \vec{f}_{21}$$

Or la propriété fondamentale liée à la troisième loi de Newton rappelée au paragraphe précédent s'écrit :

$$\vec{f}_{12} + \vec{f}_{21} = \vec{0}$$

On en conclut que la résultante des forces se réduit à la résultante des forces extérieures :

$$\vec{F} = \vec{F}_{\text{ext}} = \vec{f}_{1,\text{ext}} + \vec{f}_{2,\text{ext}}$$

1.4 Moment résultant des forces

a) Moment résultant des forces par rapport à un point

On définit le moment résultant d'un système de forces par rapport à un point O comme la somme des moments des forces par rapport à O :

$$\vec{\mathcal{M}}_O(\vec{F}) = \vec{OA}_1 \wedge \vec{f}_1 + \vec{OA}_2 \wedge \vec{f}_2$$

On établit que le moment résultant dépend du point O choisi :

$$\overrightarrow{\mathcal{M}}_{O'}(\vec{F}) = \overrightarrow{\mathcal{M}}_O(\vec{F}) + \overrightarrow{O'O} \wedge \vec{F}$$

En effet :

$$\begin{aligned} \overrightarrow{\mathcal{M}}_{O'}(\vec{F}) &= \overrightarrow{O'A_1} \wedge \vec{f}_1 + \overrightarrow{OA_2} \wedge \vec{f}_2 \\ &= (\overrightarrow{O'O} + \overrightarrow{OA_1}) \wedge \vec{f}_1 + (\overrightarrow{O'O} + \overrightarrow{OA_2}) \wedge \vec{f}_2 \\ &= \overrightarrow{O'O} \wedge (\vec{f}_1 + \vec{f}_2) + \overrightarrow{OA_1} \wedge \vec{f}_1 + \overrightarrow{OA_2} \wedge \vec{f}_2 \\ &= \overrightarrow{O'O} \wedge \vec{F} + \overrightarrow{\mathcal{M}}_O(\vec{F}) \end{aligned}$$

On a vu que la résultante des forces intérieures est nulle, on peut se demander s'il en est de même pour le moment résultant :

$$\begin{aligned} \overrightarrow{\mathcal{M}}_{O,\text{int}}(\vec{F}_{\text{int}}) &= \overrightarrow{OA_1} \wedge \vec{f}_{21} + \overrightarrow{OA_2} \wedge \vec{f}_{12} \\ &= \overrightarrow{OA_1} \wedge \vec{f}_{21} + \overrightarrow{OA_2} \wedge (-\vec{f}_{21}) \\ &= (\overrightarrow{OA_1} - \overrightarrow{OA_2}) \wedge \vec{f}_{21} \\ &= \overrightarrow{A_2A_1} \wedge \vec{f}_{21} \end{aligned}$$

en utilisant le fait que :

$$\vec{f}_{12} = -\vec{f}_{21}$$

issue de la troisième loi de Newton. Cette même loi dit que la droite d'action de $\vec{f}_{21}$ est la droite (A_2A_1) , et de ce fait, le produit vectoriel précédent est nul. Le moment résultant des forces intérieures est donc nul.

On en déduit donc que le moment résultant d'un système de forces s'exerçant sur un système de deux points matériels se réduit au moment des forces extérieures :

$$\overrightarrow{\mathcal{M}}_O(\vec{F}) = \overrightarrow{\mathcal{M}}_O(\vec{F}_{\text{ext}}) = \overrightarrow{OA_1} \wedge \vec{f}_{1,\text{ext}} + \overrightarrow{OA_2} \wedge \vec{f}_{2,\text{ext}}$$

2. Théorème de la résultante cinétique

2.1 En référentiel galiléen

Il s'agit du théorème qui remplace pour un système de deux points matériels le principe fondamental de la dynamique pour un point matériel. Il peut être désigné sous plusieurs appellations : théorème de la résultante cinétique, théorème du centre d'inertie ou théorème de la quantité de mouvement.

Il s'énonce de la manière suivante :

Dans tout référentiel galiléen, la dérivée par rapport au temps de la résultante cinétique ou quantité de mouvement du système est égale à la somme des forces extérieures agissant sur le système.

Le mouvement du centre d'inertie d'un système de deux points matériels est celui d'un point qui aurait pour masse la masse totale du système et auquel serait appliquée la résultante des forces extérieures.

On va établir ce résultat dans le cas d'un référentiel galiléen. Pour chaque point matériel A_1 ou A_2 , le principe fondamental de la dynamique donné par la deuxième loi de Newton s'écrit :

$$\begin{cases} m_1 \frac{d\vec{v}_1}{dt} = \vec{f}_1 = \vec{f}_{1,\text{ext}} + \vec{f}_{21} \\ m_2 \frac{d\vec{v}_2}{dt} = \vec{f}_2 = \vec{f}_{2,\text{ext}} + \vec{f}_{12} \end{cases}$$

avec les notations des forces précisées précédemment.

En sommant ces deux relations, on a :

$$m_1 \frac{d\vec{v}_1}{dt} + m_2 \frac{d\vec{v}_2}{dt} = \vec{f}_{1,\text{ext}} + \vec{f}_{21} + \vec{f}_{2,\text{ext}} + \vec{f}_{12}$$

soit

$$\frac{d}{dt} (m_1 \vec{v}_1 + m_2 \vec{v}_2) = \vec{f}_{1,\text{ext}} + \vec{f}_{2,\text{ext}}$$

car $\vec{f}_{12} + \vec{f}_{21} = \vec{0}$. On en déduit finalement que :

$$\frac{d\vec{P}}{dt} = \vec{F}_{\text{ext}}$$

2.2 En référentiel non galiléen

Dans le cas d'un référentiel non galiléen, le principe fondamental de la dynamique tient compte des forces d'inertie pour chaque point matériel. Or ce ne sont pas des forces intérieures : on devra en tenir compte dans le théorème de la résultante cinétique en sommant les différentes contributions correspondant à chaque point matériel. En notant $\vec{F}_{ie} = \vec{f}_{ie,1} + \vec{f}_{ie,2}$ et $\vec{F}_{ic} = \vec{f}_{ic,1} + \vec{f}_{ic,2}$ les résultantes des forces d'inertie respectivement d'entraînement et de Coriolis, on a donc :

$$\frac{d\vec{P}}{dt} = \vec{F}_{\text{ext}} + \vec{F}_{ie} + \vec{F}_{ic}$$

2.3 Application au cas d'un système de points matériels isolé

Soit un système de deux points matériels isolé. On étudie le mouvement dans un référentiel supposé galiléen. Le caractère isolé du système se traduit par le fait que la somme ou résultante des forces extérieures est nulle :

$$\vec{F}_{\text{ext}} = \vec{0}$$

Le théorème de la résultante cinétique s'écrit alors :

$$\frac{d\vec{P}}{dt} = \vec{0}$$

De plus, la masse totale du système est une constante : le caractère isolé interdit tout échange avec l'extérieur. On en déduit donc que :

$$M \frac{d\vec{v}(G)}{dt} = \vec{0}$$

donc la vitesse du centre d'inertie est constante.

Le barycentre d'un système de points matériels isolé est animé d'un mouvement rectiligne uniforme par rapport au référentiel galiléen dans lequel on étudie le mouvement.

Conséquence :

Le référentiel barycentrique est par définition en translation par rapport au référentiel galiléen à la vitesse $\vec{v}(G)$. Dans le cas d'un système isolé, $\vec{v}(G)$ est constante : le référentiel barycentrique est en translation rectiligne et uniforme par rapport au référentiel galiléen. Il est donc galiléen.

Il faut bien noter que ce n'est pas le cas en général mais que la propriété de système isolé permet d'obtenir ce résultat important.

3. Théorème du moment cinétique

3.1 Par rapport à un point fixe en référentiel galiléen ou non

La dérivée par rapport au temps du moment cinétique par rapport à un point fixe O d'un système de deux points matériels est égale au moment résultant par rapport à O de toutes les forces extérieures agissant sur le système :

$$\frac{d\vec{L}_O}{dt} = \vec{\mathcal{M}}_{O,\text{ext}}$$

Dans le cas d'un référentiel non galiléen, on tient compte des forces d'inertie d'entraînement et de Coriolis en plus des forces extérieures. Dans la suite de ce paragraphe, on considère que les forces extérieures désignent ce type de forces en référentiel galiléen et qu'elles désignent ce type de forces et les forces d'inertie en référentiel non galiléen.

Pour établir le théorème du moment cinétique par rapport à un point fixe, on procède comme pour la résultante cinétique.

Pour les points matériels A_1 et A_2 , le théorème du moment cinétique s'écrit :

$$\begin{cases} \frac{d\vec{L}_{O,1}}{dt} = \vec{M}_{O,1} = \vec{OA}_1 \wedge \vec{f}_1 = \vec{OA}_1 \wedge \vec{f}_{1,\text{ext}} + \vec{OA}_1 \wedge \vec{f}_{21} \\ \frac{d\vec{L}_{O,2}}{dt} = \vec{M}_{O,2} = \vec{OA}_2 \wedge \vec{f}_2 = \vec{OA}_2 \wedge \vec{f}_{2,\text{ext}} + \vec{OA}_2 \wedge \vec{f}_{12} \end{cases}$$

Soit, en sommant les contributions correspondant à chacun des deux points matériels du système pour obtenir le moment cinétique résultant :

$$\frac{d\vec{L}_O}{dt} = \frac{d\vec{L}_{O,1}}{dt} + \frac{d\vec{L}_{O,2}}{dt} = \vec{OA}_1 \wedge \vec{f}_{1,\text{ext}} + \vec{OA}_2 \wedge \vec{f}_{2,\text{ext}} + \vec{OA}_1 \wedge \vec{f}_{21} + \vec{OA}_2 \wedge \vec{f}_{12}$$

soit :

$$\frac{d\vec{L}_O}{dt} = \vec{OA}_1 \wedge \vec{f}_{1,\text{ext}} + \vec{OA}_2 \wedge \vec{f}_{2,\text{ext}}$$

car $\vec{OA}_1 \wedge \vec{f}_{21} + \vec{OA}_2 \wedge \vec{f}_{12} = \vec{0}$. On note que, dans le cas d'un système isolé, le moment cinétique par rapport à O est constant.

3.2 Dans le référentiel barycentrique

Le rôle particulier du centre d'inertie qui a été souligné conduit à étudier comment se traduit le théorème du moment cinétique dans le référentiel barycentrique. Soit O un point fixe du référentiel $\mathcal{R}$. Le théorème de Kœning relatif au moment cinétique s'écrit :

$$\vec{L}_{O,\mathcal{R}} = \vec{L}^* + \vec{OG} \wedge M\vec{v}_{\mathcal{R}}(G)$$

On le dérive par rapport au temps :

$$\begin{aligned} \frac{d\vec{L}_O}{dt} &= \frac{d\vec{L}^*}{dt} + \frac{d}{dt} (\vec{OG} \wedge M\vec{v}_G) = \frac{d\vec{L}^*}{dt} + \vec{v}_G \wedge M\vec{v}_G + \vec{OG} \wedge \frac{d}{dt} (M\vec{v}_G) \\ \frac{d\vec{L}_O}{dt} &= \frac{d\vec{L}^*}{dt} + \vec{OG} \wedge \vec{F}_{\text{ext}} \end{aligned}$$

Or

$$\frac{d\vec{L}_O}{dt} = \vec{M}_{O,\text{ext}} = \vec{M}_{G,\text{ext}} + \vec{OG} \wedge \vec{F}_{\text{ext}}$$

On en déduit que :

$$\overrightarrow{\mathcal{M}}_{G,\text{ext}} + \overrightarrow{OG} \wedge \overrightarrow{F}_{\text{ext}} = \frac{d\overrightarrow{L}^{\star}}{dt} + \overrightarrow{OG} \wedge \overrightarrow{F}_{\text{ext}}$$

soit

$$\frac{d\overrightarrow{L}^{\star}}{dt} = \overrightarrow{\mathcal{M}}_{G,\text{ext}}$$

Le théorème du moment cinétique par rapport à G s'écrit, dans le référentiel barycentrique, de la même manière, que celui-ci soit galiléen ou non. Les forces d'inertie n'apparaissent pas.

On peut retrouver ce résultat à partir du théorème du moment cinétique dans un référentiel non galiléen ; on détermine le moment des forces d'inertie dans le référentiel barycentrique :

$$\overrightarrow{f}_{ic} = \overrightarrow{0} \quad \text{donc} \quad \overrightarrow{\mathcal{M}}_{O,ic} = \overrightarrow{0}$$

et

$$\begin{aligned} \overrightarrow{\mathcal{M}}_{O,ic} &= \overrightarrow{GA_1} \wedge (-m_1 \overrightarrow{a}_{/\mathcal{R}}(G)) + \overrightarrow{GA_2} \wedge (-m_2 \overrightarrow{a}_{/\mathcal{R}}(G)) \\ &= (m_1 \overrightarrow{GA_1} + m_2 \overrightarrow{GA_2}) \wedge \overrightarrow{a}_{/\mathcal{R}}(G) = \overrightarrow{0} \end{aligned}$$

par définition de G .

On retrouve le résultat établi précédemment : il n'est pas nécessaire de tenir compte des forces d'inertie dans le référentiel barycentrique que celui-ci soit ou non galiléen dans l'écriture du théorème du moment cinétique en G .

4. Puissance et travail d'un système de forces

4.1 Expression

Soit deux points matériels A_1 de masse m_1 et de vitesse $\overrightarrow{v_1}$ et A_2 de masse m_2 et de vitesse $\overrightarrow{v_2}$ par rapport au référentiel d'étude $\mathcal{R}$.

La puissance par rapport à un référentiel $\mathcal{R}$ du système de forces $\overrightarrow{f_i}$ qui s'exercent sur le système de deux points s'écrit dans le référentiel $\mathcal{R}$ choisi :

$$\mathcal{P} = \overrightarrow{f_1} \cdot \overrightarrow{v_1} + \overrightarrow{f_2} \cdot \overrightarrow{v_2}$$

Il s'agit de la somme des puissances de chaque force $\overrightarrow{f_i}$ s'exerçant sur un point A_i du système considéré avec

$$\begin{cases} \overrightarrow{f_1} = \overrightarrow{f_{1,\text{ext}}} + \overrightarrow{f_{21}} \\ \overrightarrow{f_2} = \overrightarrow{f_{2,\text{ext}}} + \overrightarrow{f_{12}} \end{cases}$$

On en déduit alors :

$$\begin{aligned}\mathcal{P} &= (\vec{f}_{1,\text{ext}} + \vec{f}_{21}) \cdot \vec{v}_1 + (\vec{f}_{2,\text{ext}} + \vec{f}_{12}) \cdot \vec{v}_2 \\ &= \vec{f}_{1,\text{ext}} \cdot \vec{v}_1 + \vec{f}_{2,\text{ext}} \cdot \vec{v}_2 + \vec{f}_{12} \cdot (\vec{v}_2 - \vec{v}_1) \\ &= \mathcal{P}_{\text{ext}} + \mathcal{P}_{\text{int}}\end{aligned}$$

en notant

$$\mathcal{P}_{\text{ext}} = \vec{f}_{1,\text{ext}} \cdot \vec{v}_1 + \vec{f}_{2,\text{ext}} \cdot \vec{v}_2 \quad \text{et} \quad \mathcal{P}_{\text{int}} = \vec{f}_{12} \cdot (\vec{v}_2 - \vec{v}_1)$$

La puissance d'un système de forces s'exerçant sur un système de points matériels est donc la somme de la puissance des forces extérieures et de la puissance des forces intérieures.

Il en va de même pour le travail car $d\vec{OA}_i = \vec{v}_i dt$. Il suffit donc de multiplier toutes les équations relatives à la puissance par dt pour obtenir le résultat :

$$\delta W = \delta W_{\text{ext}} + \delta W_{\text{int}}$$

en notant

$$\delta W_{\text{ext}} = \vec{f}_{1,\text{ext}} \cdot d\vec{OA}_1 + \vec{f}_{2,\text{ext}} \cdot d\vec{OA}_2 \quad \text{et} \quad \delta W_{\text{int}} = \vec{f}_{12} \cdot (d\vec{OA}_2 - d\vec{OA}_1)$$

Le travail d'un système de forces s'exerçant sur un système de points matériels est donc la somme du travail des forces extérieures et du travail des forces intérieures.

Le travail au cours d'un déplacement s'obtient à partir du travail élémentaire en sommant les différentes contributions correspondant à chacun des déplacements élémentaires. Alors, en utilisant la linéarité de l'intégrale, on obtient la même décomposition en travail des forces intérieures et travail des forces extérieures :

$$W = \int \delta W = W_{\text{int}} + W_{\text{ext}}$$

4.2 Non nullité de la puissance et du travail des forces intérieures

Pour les deux points A_1 et A_2 , le travail élémentaire des forces intérieures est :

$$\begin{aligned}\delta W_{\text{int}} &= \mathcal{P}_{\text{int}} dt \\ &= (\vec{f}_{21} \cdot \vec{v}_1 + \vec{f}_{12} \cdot \vec{v}_2) dt \\ &= \vec{f}_{12} \cdot (\vec{v}_2 - \vec{v}_1) dt \\ &= \vec{f}_{12} \cdot (d\vec{OA}_2 - d\vec{OA}_1) \\ &= \vec{f}_{12} \cdot d\vec{A}_1 A_2\end{aligned}$$

Or $\vec{f}_{12}$ a pour droite d'action la droite $(A_1 A_2)$ et peut par conséquent s'écrire sous la forme :

$$\vec{f}_{12} = \alpha \vec{A}_1 A_2$$

On en déduit que le travail élémentaire se met sous la forme d'une différentielle de $\overrightarrow{A_1A_2}$:

$$\delta W_{\text{int}} = \alpha \overrightarrow{A_1A_2} \cdot d\overrightarrow{A_1A_2} = \frac{\alpha}{2} d(\overrightarrow{A_1A_2})^2 = \frac{\alpha}{2} d(A_1A_2)^2$$

Un travail élémentaire nul implique donc la constante de A_1A_2 .

Dans ce cas, les points A_1 et A_2 restent à égale distance au cours du déplacement. Le système formé par ces deux points est donc rigide ou indéformable lorsque le travail des forces intérieures est nul et réciproquement.

On retiendra donc que :

- le travail des forces intérieures d'un système indéformable ou rigide est nul,
- le travail des forces intérieures d'un système déformable est non nul.

Il faut bien noter que la résultante ou le moment résultant des forces intérieures sont nulles sans que cela implique que leur travail et leur puissance le soient. Il faut savoir si le système est rigide ou non.

4.3 Indépendance de la puissance et du travail des forces intérieures par rapport au référentiel

La puissance des forces intérieures s'écrit :

$$\mathcal{P}_{\text{int}} = \vec{f}_{12} \cdot \vec{v}_2 + \vec{f}_{21} \cdot \vec{v}_1 = \vec{f}_{12} \cdot (\vec{v}_2 - \vec{v}_1)$$

en tenant compte du fait que $\vec{f}_{12} = -\vec{f}_{21}$.

Dans un autre référentiel, la puissance s'exprime par :

$$\mathcal{P}'_{\text{int}} = \vec{f}_{12} \cdot \vec{v}'_2 + \vec{f}_{21} \cdot \vec{v}'_1 = \vec{f}_{12} \cdot (\vec{v}'_2 - \vec{v}'_1)$$

Or la loi de composition des vitesses s'écrit :

$$\vec{v}'_i = \vec{v}_i + \vec{v}_{ie} \quad \text{avec} \quad \vec{v}_{ie} = \vec{v}_{O'} + \vec{\Omega} \wedge \overrightarrow{OA_i}$$

Donc

$$\begin{aligned} \mathcal{P}'_{\text{int}} &= \vec{f}_{12} \cdot (\vec{v}_2 + \vec{v}_{O'} + \vec{\Omega} \wedge \overrightarrow{OA_2} - \vec{v}_1 - \vec{v}_{O'} - \vec{\Omega} \wedge \overrightarrow{OA_1}) \\ &= \vec{f}_{12} \cdot (\vec{v}_2 - \vec{v}_1 + \vec{\Omega} \wedge (\overrightarrow{OA_2} - \overrightarrow{OA_1})) \\ &= \vec{f}_{12} \cdot (\vec{v}_2 - \vec{v}_1) + \vec{f}_{12} \cdot (\vec{\Omega} \wedge \overrightarrow{A_1A_2}) \end{aligned}$$

Or $\vec{f}_{12}$ a pour droite d'action (A_1A_2) qui est perpendiculaire à $\vec{\Omega} \wedge \overrightarrow{A_1A_2}$. On en déduit : $\vec{f}_{12} \cdot (\vec{\Omega} \wedge \overrightarrow{A_1A_2}) = 0$ et

$$\mathcal{P}'_{\text{int}} = \mathcal{P}_{\text{int}}$$

La puissance des forces intérieures à un système est indépendante du référentiel choisi.

On obtient le même résultat pour le travail :

$$\delta W_{\text{int}} = \vec{f}_{12} \cdot (\overrightarrow{dOA_2} - \overrightarrow{dOA_1}) = \mathcal{P}_{\text{int}} dt$$

Donc comme $\mathcal{P}'_{\text{int}} = \mathcal{P}_{\text{int}}$, on en déduit :

$$\delta W'_{\text{int}} = \delta W_{\text{int}}$$

Le travail des forces intérieures à un système est donc indépendant du référentiel choisi.

La conséquence est d'ordre pratique : on calculera la puissance ou le travail des forces internes dans le référentiel où le calcul sera le plus facile à mener.

5. Théorème de l'énergie cinétique - Énergie mécanique

5.1 Théorème de l'énergie cinétique

a) En référentiel galiléen

Dans un référentiel galiléen, la variation d'énergie cinétique d'un système de deux points matériels entre deux instants est égale à la somme des travaux des **forces intérieures et extérieures** appliquées au système entre ces deux instants.

Le théorème de l'énergie cinétique appliqué à un point matériel s'écrit :

$$dE_c = d\left(\frac{1}{2}m_i v_i^2\right) = \delta W_i = \vec{f}_{i,\text{ext}} \cdot \vec{v}_i dt + \vec{f}_{ji} \cdot \vec{v}_i dt = \delta W_{i,\text{ext}} + \delta W_{i,\text{int}}$$

Alors, pour un système de deux points matériels, on a :

$$\begin{aligned} dE_{c1} + dE_{c2} &= \delta W_1 + \delta W_2 \\ &= \delta W_{1,\text{ext}} + \delta W_{2,\text{ext}} + \delta W_{1,\text{int}} + \delta W_{2,\text{int}} \\ &= \delta W_{\text{ext}} + \delta W_{\text{int}} \end{aligned}$$

On doit donc tenir compte dans l'application du théorème de l'énergie cinétique de TOUTES les forces, y compris les forces intérieures.

b) En référentiel non galiléen

Dans un référentiel non galiléen, la variation d'énergie cinétique d'un système de points matériels entre deux instants est égale à la somme de tous les travaux des forces appliquées au système entre ces deux instants qu'elles soient intérieures ou extérieures au système ou encore qu'il s'agisse des forces d'inertie.

Le théorème de l'énergie cinétique appliqué à un point matériel s'écrit :

$$dE_{ci} = d\left(\frac{1}{2}m_i v_i^2\right) = \delta W_i = \vec{f}_{i,\text{ext}} \cdot \vec{v}_i dt + \vec{f}_{ji} \cdot \vec{v}_i dt + \vec{f}_{i,ie} \cdot \vec{v}_i dt + \vec{f}_{i,ic} \cdot \vec{v}_i dt$$

soit

$$d\left(\frac{1}{2}m_i v_i^2\right) = \delta W_{i,\text{ext}} + \delta W_{i,\text{int}} + \delta W_{i,ie}$$

en notant $\delta W_{i,ie}$ le travail de la force d'inertie d'entraînement. On rappelle que le travail de la force de Coriolis est nul, car cette force est perpendiculaire au mouvement et ne travaille donc pas.

Alors, pour l'ensemble des deux points matériels, on a :

$$\begin{aligned} dE_{c1} + dE_{c2} &= \delta W_{1,\text{ext}} + \delta W_{1,\text{int}} + \delta W_{1,ie} + \delta W_{2,\text{ext}} + \delta W_{2,\text{int}} + \delta W_{2,ie} \\ &= \delta W_{\text{ext}} + \delta W_{\text{int}} + \delta W_{ie} \end{aligned}$$

$$\text{avec } \delta W_{ie} = \vec{f}_{1,ie} \cdot \vec{v}_1 dt + \vec{f}_{2,ie} \cdot \vec{v}_2 dt$$

Il convient donc de tenir compte dans l'application du théorème de l'énergie cinétique de TOUTES les forces, y compris les forces d'inertie et les forces intérieures.

► **Remarque :** Dans le référentiel barycentrique, le travail des forces d'inertie d'entraînement est nul. En effet :

$$\delta W_{ie}^* = -m_1 \vec{a}_{/\mathcal{R}}(G) \cdot \vec{v}_1^* - m_2 \vec{a}_{/\mathcal{R}}(G) \cdot \vec{v}_2^* = -\vec{a}_{/\mathcal{R}}(G) \cdot (m_1 \vec{v}_1^* + m_2 \vec{v}_2^*) = \vec{a}_{/\mathcal{R}}(G) \cdot \vec{P}^* = 0$$

Dans le référentiel barycentrique, le théorème de l'énergie cinétique s'écrit de la même façon que dans un référentiel galiléen, qu'il soit galiléen ou non.

5.2 Énergie potentielle

a) Énergie potentielle des forces intérieures

Dans de nombreux cas, les forces intérieures peuvent s'écrire sous la forme de forces conservatives. C'est le cas notamment des forces d'interaction gravitationnelle ou coulombienne. On définit alors une énergie potentielle.

Soit $\vec{f}_{21}$ la force exercée par le point A_2 sur le point A_1 , pour un déplacement élémentaire $d\vec{OA}_1$:

$$\delta W_{21} = \vec{f}_{21} \cdot d\vec{OA}_1$$

et la définition de l'énergie potentielle conduit à :

$$\delta W_{21} = -dEp_{21}$$

La grandeur Ep_{21} est l'énergie potentielle de la particule A_1 liée à l'action du point A_2 sur le point A_1 .

La différentielle de l'énergie potentielle des forces intérieures se met alors sous la forme :

$$dEp = dEp_{21} + dEp_{12} = -\vec{f}_{21} \cdot d\vec{OA}_1 - \vec{f}_{12} \cdot d\vec{OA}_2 = -\vec{f}_{21} \cdot d\vec{A_2A_1}$$

Ep est l'énergie potentielle d'interaction des particules entre elles.

On peut formuler quelques remarques :

1. $\vec{f}_{21}$ et $\vec{A_1A_2}$ sont colinéaires du fait du principe des actions réciproques donc

$$\vec{f}_{21} \cdot d\vec{A_1A_2} = f_{21}dA_1A_2$$

En effet :

$$\vec{f}_{21} = f_{21}\vec{u}_{12} = f_{21}\frac{\vec{A_1A_2}}{A_1A_2}$$

en notant $\vec{u}_{12}$ le vecteur unitaire de A_1A_2 dirigé de A_1 vers A_2 . Alors

$$\begin{aligned} \vec{f}_{21} \cdot d\vec{A_1A_2} &= \frac{f_{21}}{A_1A_2} \vec{A_1A_2} \cdot d\vec{A_1A_2} = \frac{f_{21}}{A_1A_2} \frac{d(A_1A_2)^2}{2} \\ &= \frac{f_{21}}{A_1A_2} (A_1A_2)dA_1A_2 = f_{21}dA_1A_2 = -dEp \end{aligned}$$

2. L'énergie d'interaction entre deux particules ne dépend que de la distance entre les points et non des variations angulaires. C'est un corollaire de la remarque précédente.
3. De la remarque précédente, on déduit que l'énergie potentielle d'interaction ne dépend que de la distance relative entre les points matériels. Elle décrit l'opposé du travail des forces intérieures du système.

Il existe une autre manière de définir l'énergie potentielle d'interaction. Si les deux particules sont infiniment éloignées l'une de l'autre, ce qu'on qualifie de particules à l'infini, il est admis que leurs interactions sont nulles ou tout au moins négligeables. De cette façon, l'énergie potentielle est telle que

$$Ep(\infty) = 0$$

Le travail qu'un opérateur extérieur doit fournir pour amener les particules de l'infini à leur position finale sans pour autant leur fournir une énergie cinétique (c'est-à-dire en effectuant l'opération infiniment lentement) s'obtient en appliquant le théorème de l'énergie cinétique dans un référentiel galiléen :

$$W_{\text{op}} + W_{\text{int}} = 0$$

puisque'il n'y a pas d'autres forces extérieures que celle de l'opérateur. On en déduit :

$$W_{\text{op}} = -W_{\text{int}} = -(Ep(\infty) - Ep) = Ep$$

Par conséquent, l'énergie potentielle d'interaction représente le travail que doit fournir un opérateur extérieur pour constituer le système à partir de particules à l'infini au sens précisé plus haut.

Exemples

Énergie potentielle d'interaction gravitationnelle

La force s'écrit :

$$\vec{f}_{21} = G \frac{m_1 m_2}{(A_1 A_2)^3} \overrightarrow{A_1 A_2}$$

Donc

$$-dEp = \vec{f}_{21} \cdot d\overrightarrow{A_2 A_1} = f_{21} dA_1 A_2 = -G \frac{m_1 m_2}{(A_1 A_2)^2} dA_1 A_2 = G m_1 m_2 d\left(\frac{1}{A_1 A_2}\right)$$

On en déduit :

$$Ep = -G \frac{m_1 m_2}{A_1 A_2} + C$$

Or $Ep(\infty) = 0$ donc $C = 0$ et

$$Ep = -G \frac{m_1 m_2}{A_1 A_2}$$

Énergie potentielle d'interaction coulombienne

La force s'écrit :

$$\vec{f}_{21} = -\frac{1}{4\pi\epsilon_0} \frac{q_1 q_2}{(A_1 A_2)^3} \overrightarrow{A_1 A_2}$$

Donc, par une démonstration analogue à la précédente (obtenue en remplaçant $-Gm_1 m_2$ par $\frac{q_1 q_2}{4\pi\epsilon_0}$), on a :

$$Ep = \frac{1}{4\pi\epsilon_0} \frac{q_1 q_2}{A_1 A_2} = q_1 V_2(A_1)$$

en notant

$$V_2(A_1) = \frac{1}{4\pi\epsilon_0} \frac{q_2}{A_1 A_2}$$

le potentiel électrique créé au point A_1 par la charge q_2 située en A_2 (Cf. cours d'électromagnétisme).

b) Énergie potentielle des forces extérieures

On ne peut définir une énergie potentielle que pour les forces extérieures dérivant d'un potentiel ou forces conservatives. Il convient alors de sommer l'ensemble des contributions s'appliquant sur chaque point matériel constituant le système.

Énergie potentielle de pesanteur

Le travail s'écrit :

$$\delta W_i = (m_i \vec{g}) \cdot d\vec{OA}_i = d(m_i \vec{g} \cdot \vec{OA}_i)$$

en considérant que le champ de pesanteur est uniforme.

On en déduit :

$$\begin{aligned} \delta W &= d(m_1 \vec{g} \cdot \vec{OA}_1) + d(m_2 \vec{g} \cdot \vec{OA}_2) \\ &= d(\vec{g} \cdot (m_1 \vec{OA}_1 + m_2 \vec{OA}_2)) \\ &= d(\vec{g} \cdot M\vec{OG}) \end{aligned}$$

par définition du centre d'inertie. On en déduit :

$$Ep = Mgz + C$$

où z est l'**altitude du centre d'inertie** (avec l'axe orienté vers le haut). On trouve la même expression que pour un point matériel en considérant que le centre d'inertie est doté de la masse totale.

Énergie potentielle associée aux forces d'inertie d'entraînement

De la même manière, on montre que l'énergie potentielle de la force centrifuge $m_i \Omega^2 \vec{H}_i \vec{A}_i$ en notant H_i la projection de A_i sur l'axe de rotation vaut :

$$Ep = -\frac{1}{2} (m_1 (H_1 A_1)^2 + m_2 (H_2 A_2)^2) \Omega^2 + C$$

en notant C la constante globale.

5.3 Énergie mécanique

a) Théorème de l'énergie mécanique

On distingue les forces conservatives de celles qui ne le sont pas, et ce, aussi bien pour les forces intérieures que pour les forces extérieures.

$$\delta W_{\text{int}} = \delta W'_{\text{int}} - dEp_{\text{int}}$$

$$\delta W_{\text{ext}} = \delta W'_{\text{ext}} - dEp_{\text{ext}}$$

en notant $\delta W'_{\text{int}}$ et $\delta W'_{\text{ext}}$ les travaux des forces non conservatives respectivement intérieures et extérieures. Les travaux des forces conservatives se mettent sous la forme de l'opposé d'une variation d'énergie potentielle. Le théorème de l'énergie cinétique s'écrit alors :

$$dE_c = \delta W'_{\text{int}} - dE_{p_{\text{int}}} + \delta W'_{\text{ext}} - dE_{p_{\text{ext}}}$$

soit

$$d(E_c + E_{p_{\text{int}}} + E_{p_{\text{ext}}}) = \delta W'_{\text{int}} + \delta W'_{\text{ext}}$$

On obtient ainsi le théorème de l'énergie mécanique en notant $E_m = E_c + E_{p_{\text{int}}} + E_{p_{\text{ext}}}$:

$$dE_m = \delta W'_{\text{int}} + \delta W'_{\text{ext}}$$

La variation d'énergie mécanique est égale à la somme des travaux des forces intérieures et extérieures non conservatives, c'est-à-dire de l'ensemble des forces ne dérivant pas d'un potentiel.

b) Cas d'un système isolé ou pseudo-isolé

Un système est dit isolé ou pseudo-isolé lorsqu'il n'est soumis à aucune force extérieure ou que l'influence des forces extérieures est nulle.

Dans ce cas, on a : $\delta W'_{\text{ext}} = dE_{p_{\text{ext}}} = 0$ et le théorème de l'énergie mécanique devient

$$d(E_c + E_{p_{\text{int}}}) = \delta W'_{\text{int}}$$

L'énergie mécanique d'un tel système $E_m = E_{p_{\text{int}}} + E_c$ n'est alors en général pas conservée. Ce ne sera le cas que si

$$\delta W'_{\text{int}} = 0$$

L'énergie mécanique n'est donc pas une grandeur conservative, sa conservation n'est qu'occasionnelle.

On introduira en thermodynamique un nouveau concept :

- l'énergie interne qui permettra d'élaborer une grandeur énergétique conservative ;
- l'énergie totale, somme de l'énergie mécanique et de l'énergie interne.

On verra qu'alors on préfère souvent inclure le terme d'énergie potentielle des forces intérieures dans l'énergie interne plutôt que dans l'énergie mécanique. Il ne s'agit que d'une définition qui ne change rien au phénomène physique de conservation de l'énergie totale pour un système isolé.

A. Application directe du cours

1. Conservation de l'énergie mécanique - Vrai ou faux

Les affirmations suivantes sont-elles vraies ou fausses ?

L'énergie mécanique d'un système de points matériels est conservée :

1. si le système est isolé,
2. si le travail des forces ne dépend pas du chemin suivi,
3. si les forces dérivent d'un potentiel,
4. si l'énergie cinétique est constante.

2. Deux masses liées par un ressort

Deux masses m_1 et m_2 sont fixées aux extrémités M_1 et M_2 d'un ressort de raideur k et de longueur à vide ℓ_0 , de masse négligeable. Elles sont assujetties à se déplacer sans frottement sur un axe horizontal Ox .

À $t = 0$, le système est au repos ; une particule de masse m se déplaçant sur le plan horizontal entre en collision avec la masse M_1 ; juste après le choc, la vitesse de M_1 est v_0 colinéaire au ressort, la masse m est éjectée vers l'arrière.

1. Déterminer après le choc le mouvement du centre d'inertie G des deux masses M_1 et M_2 .
2. Déterminer la loi de variation $\ell(t)$ de la longueur du ressort de 2 manières différentes :
 - dans le référentiel barycentrique, en introduisant la particule fictive ;
 - directement dans le référentiel absolu en raisonnant sur chaque masse.

B. Exercices et problèmes

1. Deux masses sur un plan incliné

Une masse m_1 est mobile sur un plan horizontal avec un coefficient de frottement cinétique f . Elle est reliée par un fil passant sur une poulie idéale à une masse m_2 située initialement à une hauteur h au-dessus du sol. On lâche le système à l'instant $t = 0$. La masse m_1 parcourt une distance $h + d$ avant de s'arrêter. Montrer que ces informations permettent de calculer la valeur du coefficient de frottement f en fonction de m_1 , m_2 , h et d . On rappelle les lois de Coulomb relatives au frottement : la réaction du support $\vec{R}$ se décompose en une composante tangentielle $\vec{R}_T$ et une composante normale $\vec{R}_N$ telles que $\vec{R} = \vec{R}_T + \vec{R}_N$ et $R_T = fR_N$ quand il y a glissement.

2. Pendule sur un chariot, d'après Centrale 1995

On étudie l'expérience amusante suivante : un chariot, sur lequel est fixé un pendule pesant, peut rouler entre deux butoirs. On considère que le chariot est assimilable à un point matériel C , de masse M et que son déplacement suivant l'axe Ox se fait sans frottement (figure 29.1) grâce à un système type "coussin d'air".

Le pendule pourra être considéré comme un pendule simple (point B) de masse m et de longueur $AB = L$. Tout le système de fixation du pendule (potence, fil, ...) est de masse négligeable.

Dans tout l'exercice, on raisonne dans le référentiel galiléen lié au sol.

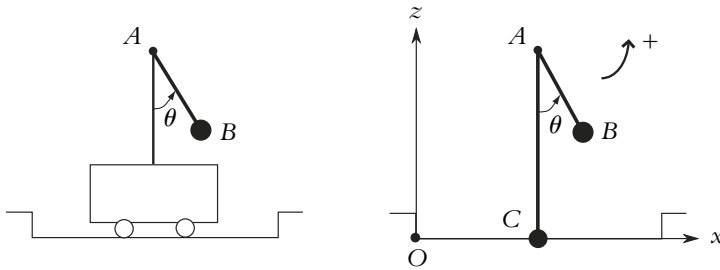


Figure 29.1

1. On suppose que le chariot n'est pas en butée.

a) On écarte le pendule de sa position d'équilibre d'un angle $\theta_0 \in [-\frac{\pi}{2}, 0]$ et on lâche B , sans communiquer de vitesse initiale au système. Établir la relation entre $\dot{x}$ vitesse du chariot et $\dot{\theta}$ vitesse angulaire du pendule. On posera $\alpha = \frac{m}{m + M}$.

b) Que peut-on dire du mouvement du centre de gravité des deux points ?

On change l'origine du repère ; la nouvelle origine O' étant sur l'axe Ox , à la verticale de G . On note (X_B, Z_B) les coordonnées de B et (X_C, Z_C) celles de C dans ce nouveau repère.

c) Établir une relation entre X_B et X_C . Décrire qualitativement le mouvement du chariot en fonction de la position de B .

Déterminer l'expression de la cote de G en fonction de θ , L , α et $h = CA$.

d) Montrer que les coordonnées de B sont :

$$X_B = (1 - \alpha)L \sin \theta \quad \text{et} \quad Z_B = (1 - \alpha)(h - L \cos \theta)$$

En déduire la trajectoire de B .

2. On procède de même que précédemment, mais cette fois-ci le chariot est contre la butée de gauche.

À partir de quel instant le système est-il pseudo-isolé (en ce qui concerne le mouvement horizontal) ? À ce même instant quelle est la différence avec la situation précédente ? En déduire la nouvelle relation entre $\dot{x}$ et $\dot{\theta}$.

Dans la suite on considère le chariot loin des butées.

3. a) Établir l'expression de l'énergie cinétique totale en fonction de θ et de ses dérivées.

b) Établir l'expression de l'énergie potentielle totale, en fonction de θ .

c) Exprimer l'énergie mécanique totale. Quelle propriété a-t-elle ? En déduire l'équation différentielle du mouvement. Simplifier cette équation en ne considérant que les petites oscillations et en déduire la période correspondante.

3. Ressort en rotation

Deux points matériels A et B de même masse m sont reliés par un ressort sans masse de longueur à vide ℓ_0 et de raideur k . Le point A est également relié à un point fixe O par l'intermédiaire d'un fil inextensible de longueur L . L'ensemble tourne sans frottement dans un plan horizontal à la vitesse angulaire ω constante autour du point fixe O . Pendant cette première phase la distance $r = AB$ reste constante.

On suppose que le ressort ne se tord pas et que $m\omega^2 < k$.

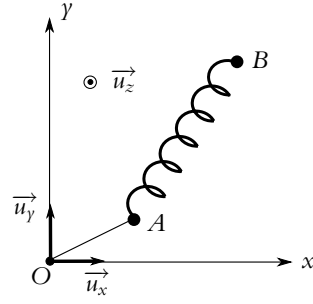


Figure 29.2

1. Montrer que les points O , A et B sont alignés. Calculer la tension T_f du fil et la longueur ℓ_1 du ressort. Application : $k = 2m\omega^2$; exprimer ℓ_1 en fonction de L et ℓ_0 .
2. À $t = 0$, on brûle le fil entre O et A . Les points A et B se trouvent à cet instant sur l'axe Ox .
 - a) Déterminer le mouvement ultérieur du centre d'inertie G des deux points matériels. Calculer sa vitesse et sa position. Application : $k = 2m\omega^2$; exprimer la vitesse V_G en fonction de L , ℓ_1 et ω .
 - b) Définir le référentiel barycentrique. Qu'a-t-il de remarquable ?

La suite du problème sera traitée dans le référentiel barycentrique.

3. a) Déterminer les expressions des positions de A et B par rapport à G et de leurs quantités de mouvement en fonction de m et de $\overrightarrow{AB}$.
- b) Calculer l'énergie cinétique totale E_c^* , le moment cinétique total σ^* et l'énergie potentielle d'interaction E_p du système (A, B) ; en déduire la position par rapport à G et la masse d'un point matériel M qui aurait cette énergie cinétique, ce moment cinétique et cette énergie potentielle.
- c) Calculer la vitesse initiale V_1 de M .
Application : $k = 2m\omega^2$; exprimer V_1 en fonction de L , ℓ_0 et ω .
- d) Quelles sont les forces subies par M ?
- e) En déduire les deux grandeurs constantes au cours du mouvement et calculer leurs valeurs en fonction de ℓ_1 , ℓ_0 , k , m et ω .
Application : $k = 2m\omega^2$; les exprimer ensuite en fonction de m , L , ℓ_0 et ω .
- f) Définir le potentiel effectif $E_{p,\text{eff}}$; donner son expression pour $k = 2m\omega^2$ en fonction de m , ω , L , ℓ_0 et $r = AB$; tracer son allure. On montrera qu'il présente un minimum pour $r_m < \ell_1$, pour cela on pourra calculer $\frac{dE_{p,\text{eff}}}{dr}(\ell_1)$.
- g) Positionner sur la courbe précédente le point correspondant à $t = 0$. Déterminer la longueur maximale du ressort au cours du temps.
- h) Décrire avec précision le mouvement de A et B dans le référentiel barycentrique, puis dans le référentiel de départ.

30

Systèmes de deux points matériels isolés

Ce chapitre s'intéresse au système de points matériels le plus simple : celui de deux points isolés en interaction. On peut considérer qu'il s'agit par exemple de particules assimilées à des masses ponctuelles. Ces deux points matériels exercent l'un sur l'autre une force d'interaction.

1. Décomposition du mouvement

Dans ce paragraphe, on ne se préoccupe pas de la nature de l'interaction et on ne formule aucune hypothèse sur la force correspondante.

1.1 Mouvement du centre d'inertie

Soient deux points matériels A_1 et A_2 de masse respective m_1 et m_2 et animés d'une vitesse respective $\vec{v}_1$ et $\vec{v}_2$ dans le référentiel d'étude $\mathcal{R}$.

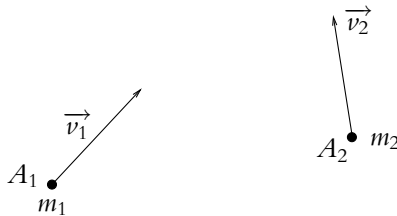


Figure 30.1 Mouvement relatif de deux points isolés.

En considérant le système dans son ensemble, on cherche le mouvement du centre d'inertie G du système défini comme le barycentre de A_1 affecté du poids m_1 et de A_2 affecté du poids m_2 . Ce mouvement est obtenu, d'après le théorème de la résultante cinétique appliqué au système dans son ensemble en étudiant le mouvement du point

matériel G auquel on affecte la totalité de la masse du système $m_1 + m_2$. Il est soumis à l'ensemble des actions extérieures au système. On ne tient pas compte de l'interaction entre les deux points matériels pour déterminer la trajectoire du point G .

1.2 Mouvement relatif

À ce mouvement du centre d'inertie, il faut ajouter le mouvement relatif des deux points matériels entre eux.

Le mouvement relatif du point A_2 par rapport au point A_1 est défini par les relations suivantes :

- la position relative :

$$\vec{r} = \overrightarrow{A_1 A_2} = \overrightarrow{OA_2} - \overrightarrow{OA_1}$$

en notant O un point fixe dans le référentiel considéré,

- la vitesse relative :

$$\vec{v} = \vec{v}_2 - \vec{v}_1$$

et

$$\vec{v}_2 - \vec{v}_1 = \frac{d\overrightarrow{OA_2}}{dt} - \frac{d\overrightarrow{OA_1}}{dt} = \frac{d(\overrightarrow{OA_2} - \overrightarrow{OA_1})}{dt} = \frac{d\overrightarrow{A_1 A_2}}{dt} = \frac{d\vec{r}}{dt}$$

Il s'agit bien d'une vitesse relative d'après la loi de composition des vitesses. En effet, en considérant le référentiel $\mathcal{R}_1$ lié au point A_1 et en translation par rapport au référentiel fixe $\mathcal{R}$, la loi de composition des vitesses s'écrit, pour le point A_2 :

$$\vec{v}_2 = \vec{v}_{/\mathcal{R}}(A_2) = \vec{v}_e + \vec{v}_{/\mathcal{R}_1}(A_2)$$

avec la vitesse d'entraînement $\vec{v}_e = \vec{v}_{/\mathcal{R}}(A_1) = \vec{v}_1$.

- l'accélération relative :

$$\vec{a} = \vec{a}_2 - \vec{a}_1 = \frac{d\vec{v}_2}{dt} - \frac{d\vec{v}_1}{dt} = \frac{d(\vec{v}_2 - \vec{v}_1)}{dt} = \frac{d\vec{v}}{dt}$$

2. Référentiel barycentrique et éléments dynamiques

2.1 Référentiel barycentrique

Soit O un point fixe dans le référentiel fixe $\mathcal{R}$.

Par définition du centre d'inertie G :

$$\overrightarrow{OG} = \frac{1}{m_1 + m_2} (m_1 \overrightarrow{OA_1} + m_2 \overrightarrow{OA_2})$$

On rappelle que le référentiel barycentrique noté $\mathcal{R}^*$ est le référentiel dont l'origine est le centre d'inertie G et qui est animé d'un mouvement de translation à la vitesse du centre d'inertie $\vec{v}(G)$.

Cette vitesse s'obtient à partir de la relation précédente :

$$\vec{v}(G) = \frac{d\vec{OG}}{dt} = \frac{1}{m_1 + m_2} \left(m_1 \frac{d\vec{OA}_1}{dt} + m_2 \frac{d\vec{OA}_2}{dt} \right) = \frac{m_1 \vec{v}_1 + m_2 \vec{v}_2}{m_1 + m_2}$$

La quantité de mouvement totale du système s'exprime en fonction de la vitesse du centre d'inertie comme cela a été vu au chapitre 28 :

$$\vec{P} = m_1 \vec{v}_1 + m_2 \vec{v}_2 = (m_1 + m_2) \vec{v}(G)$$

2.2 Expression des quantités de mouvement dans $\mathcal{R}^*$

La quantité de mouvement totale du système dans $\mathcal{R}^*$ est nulle : par définition même du référentiel barycentrique, la vitesse du centre d'inertie G dans ce référentiel est nulle. Donc

$$\vec{P}^* = (m_1 + m_2) \vec{v}^*(G) = \vec{0}$$

On peut alors expliciter les quantités de mouvement de chaque point matériel dans $\mathcal{R}^*$. On note tout d'abord qu'elles sont opposées l'une de l'autre :

$$\vec{P}^* = \vec{p}_1^* + \vec{p}_2^* = \vec{0}$$

d'où :

$$\vec{p}_1^* = -\vec{p}_2^*$$

D'autre part, on peut expliciter la loi de composition des vitesses dans le référentiel barycentrique :

$$\vec{v}_1 = \vec{v}(G) + \vec{v}_1^*$$

donc

$$\vec{v}_1^* = \vec{v}_1 - \vec{v}(G)$$

De même, on a :

$$\vec{v}_2^* = \vec{v}_2 - \vec{v}(G)$$

Comme

$$\vec{v}_G = \frac{m_1 \vec{v}_1 + m_2 \vec{v}_2}{m_1 + m_2}$$

on en déduit :

$$\vec{v}_1^* = -\frac{m_2}{m_1 + m_2} \vec{v}$$

avec $\vec{v} = \vec{v}_2 - \vec{v}_1$ la vitesse relative de A_2 par rapport à A_1 .

De même, on a :

$$\vec{v}_2^* = \frac{m_1}{m_1 + m_2} \vec{v}$$

On obtient alors les expressions des quantités de mouvement dans le référentiel barycentrique :

$$\vec{p}_1^* = m_1 \vec{v}_1^* = -\frac{m_1 m_2}{m_1 + m_2} (\vec{v}_2 - \vec{v}_1)$$

et

$$\vec{p}_2^* = m_2 \vec{v}_2^* = \frac{m_1 m_2}{m_1 + m_2} (\vec{v}_2 - \vec{v}_1)$$

2.3 Masse réduite

Dans les expressions qui viennent d'être obtenues apparaît une quantité homogène à une masse :

$$\mu = \frac{m_1 m_2}{m_1 + m_2}$$

On l'appelle *masse réduite* du système. Son introduction permet d'écrire les quantités de mouvement précédentes sous la forme :

$$\vec{p}_2^* = -\vec{p}_1^* = \mu \vec{v} \quad (30.1)$$

en fonction de la vitesse relative $\vec{v}$ de A_2 par rapport à A_1 .

2.4 Expression du moment cinétique du système dans $\mathcal{R}^*$

$$\vec{L}^* = \overrightarrow{GA_1} \wedge \vec{p}_1^* + \overrightarrow{GA_2} \wedge \vec{p}_2^*$$

Or d'après la relation (30.1) :

$$\vec{p}_2^* = -\vec{p}_1^* = \mu \vec{v}$$

D'où :

$$\vec{L}^* = (\overrightarrow{GA_2} - \overrightarrow{GA_1}) \wedge \mu \vec{v} = -\overrightarrow{A_1 A_2} \wedge \mu \vec{v}$$

soit

$$\vec{L}^* = \overrightarrow{A_1 A_2} \wedge \mu \vec{v}$$

2.5 Expression de l'énergie cinétique du système dans $\mathcal{R}^*$

Dans le référentiel barycentrique, l'énergie cinétique s'écrit :

$$Ec^* = \frac{1}{2} m_1 v_1^{*2} + \frac{1}{2} m_2 v_2^{*2} = \frac{p_1^{*2}}{2m_1} + \frac{p_2^{*2}}{2m_2} = \left(\frac{1}{m_1} + \frac{1}{m_2} \right) \frac{\mu^2 v^2}{2} = \frac{m_1 + m_2}{m_1 m_2} \frac{\mu^2 v^2}{2}$$

soit

$$Ec^* = \frac{1}{2} \mu v^2$$

► **Remarque** : On notera que l'expression de l'énergie cinétique sous la forme $\frac{p^2}{2m}$ au lieu de $\frac{1}{2}mv^2$ permet d'obtenir rapidement le résultat puisque $p_1^{*2} = p_2^{*2}$.

3. Particule fictive

3.1 Définition

L'ensemble des expressions qui viennent d'être obtenues suggère d'introduire une particule fictive appelée aussi mobile réduit ou mobile équivalent. En effet, on voit apparaître des expressions analogues à celles d'un point matériel A dont les caractéristiques seraient les suivantes :

- une masse égale à la masse réduite :

$$\mu = \frac{m_1 m_2}{m_1 + m_2}$$

- une position définie par le vecteur :

$$\vec{r} = \overrightarrow{GA} = \overrightarrow{A_1 A_2}$$

3.2 Éléments dynamiques

Cette particule dite fictive permet de fournir une interprétation aux expressions de la quantité de mouvement, du moment cinétique et de l'énergie cinétique du système dans le référentiel barycentrique :

- quantité de mouvement :

$$\vec{p}_2^* = \mu \vec{v} = \mu \frac{d\vec{r}}{dt} \quad (\text{cf. paragraphe 1.2})$$

La quantité de mouvement de la particule fictive est égale à la quantité de mouvement du point matériel A_2 dans le référentiel barycentrique.

- moment cinétique :

$$\vec{L}^* = \overrightarrow{A_1 A_2} \wedge \mu \vec{v} = \vec{r} \wedge \mu \vec{v}$$

Le moment cinétique de la particule fictive est celui du système total exprimé dans le référentiel barycentrique.

- énergie cinétique :

$$Ec^* = \frac{1}{2} \mu v^2$$

L'énergie cinétique de la particule fictive est celle du système total exprimée dans le référentiel barycentrique.

3.3 Conséquences

On déduit de ce qui précède que l'étude d'un système de deux points matériels A_1 et A_2 en interaction se ramène à :

1. l'étude du mouvement de son centre d'inertie G soumis aux actions extérieures au système,
2. l'étude dans le référentiel barycentrique d'une particule fictive A de masse μ telle que :

$$\left\{ \begin{array}{l} \vec{r} = \overrightarrow{GA} = \overrightarrow{A_1A_2} \\ \mu = \frac{m_1 m_2}{m_1 + m_2} \\ \vec{p}^* = \vec{p}_2^* = -\vec{p}_1^* \\ \vec{L}^* = \vec{r} \wedge \vec{p}^* \\ Ec^* = \frac{1}{2} \mu v^{*2} \end{array} \right.$$

Le mouvement des points matériels A_1 et A_2 peut être déduit par homothétie de la trajectoire du mobile fictif. En effet, on a les relations suivantes :

$$\left\{ \begin{array}{l} m_1 \overrightarrow{GA_1} + m_2 \overrightarrow{GA_2} = \vec{0} \\ \overrightarrow{GA} = \overrightarrow{A_1A_2} = \overrightarrow{GA_2} - \overrightarrow{GA_1} \end{array} \right.$$

La résolution de ce système de deux équations à deux inconnues $\overrightarrow{GA_1}$ et $\overrightarrow{GA_2}$ donne :

$$\left\{ \begin{array}{l} \overrightarrow{GA_1} = -\frac{m_2}{m_1 + m_2} \overrightarrow{GA} = -\frac{\mu}{m_1} \overrightarrow{GA} \\ \overrightarrow{GA_2} = \frac{m_1}{m_1 + m_2} \overrightarrow{GA} = \frac{\mu}{m_2} \overrightarrow{GA} \end{array} \right.$$

Les trajectoires des points matériels A_1 et A_2 ainsi que celle du mobile fictif A sont homothétiques les unes des autres dans le référentiel barycentrique.

4. Lois de conservation et conséquences

4.1 Hypothèse d'un système isolé

Jusqu'ici on n'a formulé aucune hypothèse sur le système si ce n'est qu'il est composé de deux points matériels. On se place désormais dans le cas particulier d'un système isolé. Cela signifie qu'il n'y a pas de forces extérieures.

On étudie le mouvement dans un référentiel supposé galiléen.

4.2 Caractère galiléen du référentiel barycentrique

Dans le cadre de cette hypothèse d'un système isolé, le théorème de la résultante cinétique s'écrit :

$$(m_1 + m_2) \frac{d\vec{v}(G)}{dt} = \vec{0}$$

puisque la somme des forces extérieures s'appliquant sur le système est nulle et que ce sont les seules forces qui interviennent dans ce théorème.

On en déduit que la vitesse du centre d'inertie est une constante :

$$\vec{v}(G) = \overrightarrow{cste}$$

Le centre d'inertie est donc animé d'un mouvement de translation rectiligne et uniforme.

Il en résulte que le référentiel barycentrique, défini comme un référentiel en translation à la vitesse $\vec{v}(G)$, est un référentiel galiléen par définition même de tels référentiels.



Il ne faut pas en déduire que tout référentiel barycentrique est galiléen, ce n'est le cas que lorsque le système est isolé comme ici.

4.3 Trajectoires

Du fait du caractère galiléen du référentiel barycentrique pour les systèmes isolés, on peut écrire le principe fondamental de la dynamique pour chacune des masses dans le référentiel barycentrique :

$$\frac{d\vec{p}_1^*}{dt} = \overrightarrow{f_{2 \rightarrow 1}} = -\mu \frac{d\vec{v}}{dt}$$

et

$$\frac{d\vec{p}_2^*}{dt} = \overrightarrow{f_{1 \rightarrow 2}} = \mu \frac{d\vec{v}}{dt}$$

Le principe des actions réciproques est bien compatible avec ces deux équations :

$$\overrightarrow{f_{2 \rightarrow 1}} = -\overrightarrow{f_{1 \rightarrow 2}} = -\mu \frac{d\vec{v}}{dt}$$

De plus, tout se passe comme si la particule fictive A était soumise à la force $\overrightarrow{f_{1 \rightarrow 2}}$. On note qu'il s'agit d'une *force centrale* telle qu'on la définira au chapitre suivant, c'est-à-dire une force dont la direction passe par un point fixe (ici, G).

L'intégration de cette équation différentielle permet de déterminer la vitesse relative $\vec{v}$ et la position de la particule fictive $\vec{r}$ en tenant compte des conditions initiales.

Cette équation peut donc être considérée soit comme l'équation du mouvement permettant d'obtenir le mouvement relatif des particules, soit comme l'équation du mouvement de la particule fictive de masse μ et de position $\overrightarrow{GA} = \overrightarrow{A_1A_2}$ dans le référentiel barycentrique.

On rappelle que les trajectoires de points A_1 et A_2 se déduisent par homothétie de la trajectoire de la particule fictive (Cf. paragraphe 3.3).

4.4 Conservation du moment cinétique

Le moment cinétique du système est constant dans le référentiel barycentrique.

En effet, le théorème du moment cinétique appliqué à l'ensemble des deux particules dans le référentiel barycentrique donne :

$$\frac{d\vec{L}^*}{dt} = \overrightarrow{\mathcal{M}}_{G,\text{ext}}$$

Comme il n'y a pas de forces extérieures, on a :

$$\overrightarrow{\mathcal{M}}_{G,\text{ext}} = \vec{0}$$

et

$$\frac{d\vec{L}^*}{dt} = \vec{0}$$

ce qui permet de déduire que le moment cinétique du système dans le référentiel barycentrique est constant :

$$\vec{L}^* = \text{cste}$$

4.5 Conséquences de la conservation du moment cinétique

Le moment cinétique du système est constant dans le référentiel barycentrique, on peut l'écrire sous la forme :

$$\vec{L}^* = \vec{r}_0 \wedge \mu \vec{v}_0 = \mu \vec{C}$$

en notant $\vec{C} = \vec{r}_0 \wedge \vec{v}_0$.

On doit distinguer deux cas suivant qu'il s'agit d'un vecteur nul ou non.

- $\vec{C} = \vec{0}$:
Dans ce cas, $\vec{r}_0$ et $\vec{v}_0$ sont colinéaires. Le mouvement s'effectue sur la droite dirigée par ces vecteurs et passant par la position initiale.
- $\vec{C} \neq \vec{0}$:
Le mouvement s'effectue perpendiculairement à une direction fixe donnée par le vecteur $\vec{C}$. On aura donc un mouvement plan dans le plan défini par $(G, \vec{r}_0, \vec{v}_0)$.

En utilisant les coordonnées polaires dans ce plan, on obtient :

$$\begin{cases} \vec{r} = r \vec{u}_r \\ \vec{v} = \dot{r} \vec{u}_r + r \dot{\theta} \vec{u}_\theta \end{cases}$$

$\vec{C}$ s'écrit alors :

$$\vec{C} = r \vec{u}_r \wedge (\dot{r} \vec{u}_r + r \dot{\theta} \vec{u}_\theta) = r^2 \dot{\theta} \vec{u}_z$$

On en déduit que la grandeur

$$C = r^2 \dot{\theta}$$

est une constante du mouvement. C est appelée *constante des aires* et vaut deux fois la vitesse aréolaire ou aire balayée par le rayon-vecteur $\vec{GA}$ pendant l'intervalle de temps dt :

$$\frac{d\mathcal{A}}{dt} = \frac{r^2 \dot{\theta}}{2}$$

Ce point sera revu dans le chapitre sur les forces centrales.

On peut noter également que $\dot{\theta}$ garde un signe constant au cours du temps : la particule fictive tourne toujours dans le même sens.

On retiendra que la conservation du moment cinétique traduit le fait que le mouvement est plan et que le rayon-vecteur balaie une aire constante par unité de temps.

4.6 Forces conservatives et conservation de l'énergie

On formule maintenant une hypothèse supplémentaire à celle d'un système isolé de deux points matériels : on suppose que les forces d'interaction entre les deux points matériels sont conservatives et donc dérivent d'un potentiel.

Dans ce cas, le théorème de l'énergie mécanique :

$$d(Ec + Ep_{\text{int}} + Ep_{\text{ext}}) = \delta W'_{\text{int}} + \delta W'_{\text{ext}}$$

s'écrit :

$$d(Ec + Ep_{\text{int}}) = 0$$

L'énergie mécanique $Em = Ec + Ep_{\text{int}}$ est donc ici une constante du mouvement.

On utilise les coordonnées polaires dans le plan du mouvement. L'énergie cinétique s'écrit :

$$Ec = \frac{1}{2} \mu \left(\dot{r}^2 + (r \dot{\theta})^2 \right)$$

soit en introduisant la constante des aires $C = r^2 \dot{\theta}$:

$$E_c = \frac{1}{2} \mu \left(\dot{r}^2 + \frac{C^2}{r^2} \right)$$

L'expression de l'énergie mécanique devient :

$$E_m = E_c + E_{p_{\text{int}}}(r) = \frac{1}{2} \mu \dot{r}^2 + \frac{\mu C^2}{2r^2} + E_{p_{\text{int}}}(r)$$

► **Remarque** : On peut également se placer dans le référentiel barycentrique. D'après le deuxième théorème de Koenig, l'énergie cinétique s'écrit :

$$E_c = E_c^* + \frac{1}{2} M v^2(G)$$

On en déduit l'expression de l'énergie mécanique :

$$E_m = E_m^* + \frac{1}{2} M v^2(G)$$

Comme le système est isolé, $\frac{1}{2} M v^2(G)$ est une constante et l'énergie mécanique exprimée dans le référentiel barycentrique est constante.

5. Analyse qualitative du mouvement

Dans ce paragraphe, on se place dans l'hypothèse d'un système isolé de deux points matériels dont l'interaction est conservative. Tous les résultats précédents sont applicables.

5.1 Énergie potentielle effective

L'expression obtenue pour la conservation de l'énergie mécanique conduit à introduire la notion d'*énergie potentielle effective ou efficace* :

$$U_{\text{eff}}(r) = E_{p_{\text{int}}}(r) + \frac{\mu C^2}{2} \frac{1}{r^2}$$

On se ramène ainsi à un problème unidimensionnel en r . Il suffit de résoudre cette équation de conservation de l'énergie mécanique d'une particule fictive de masse μ placée dans un potentiel effectif U_{eff} .

Dans la suite de ce paragraphe, on va discuter la nature du mouvement suivant l'allure de ce potentiel efficace et les valeurs de la constante de l'énergie mécanique E_0 .

5.2 Puits d'énergie potentielle effective

Dans ce cas, la courbe donnant l'énergie potentielle effective en fonction de r passe par un minimum.

Sa valeur limite quand $r \rightarrow +\infty$ est zéro ; en effet, les interactions sont nulles lorsque les deux points matériels sont infiniment éloignés l'un de l'autre (donc quand $r \rightarrow +\infty$, l'énergie potentielle et par conséquent l'énergie potentielle efficace tendent vers 0) et $\lim_{r \rightarrow +\infty} \frac{\mu C^2}{2r^2} = 0$. On obtient par exemple l'allure suivante (cette courbe n'est qu'un exemple destiné à illustrer le raisonnement) :

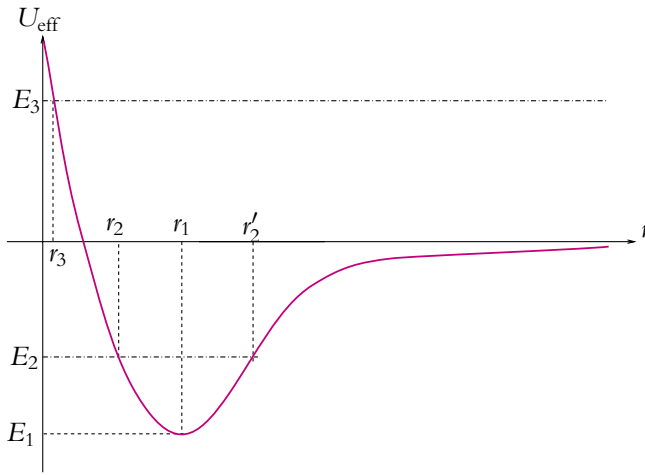


Figure 30.2 Puits d'énergie potentielle effective.

Le terme $\frac{1}{2}\mu\dot{r}^2$ est positif par définition, on en déduit que

$$E_0 - U_{\text{eff}}(r) \geq 0$$

On distingue donc plusieurs types de mouvement suivant la valeur de E_0 :

- $E_0 \geq 0$ (par exemple $E_0 = E_3$ sur le schéma ci-dessus) : l'inégalité précédente est vérifiée si $r \geq r_3$ d'après le graphe. On a alors un *état libre* où la particule fictive peut atteindre l'infini.
- $E_1 \leq E_0 \leq 0$ (par exemple $E_0 = E_2$ sur le schéma ci-dessus) : l'inégalité précédente est vérifiée si $r_2 \leq r \leq r'_2$ d'après le graphe. On a alors un *état lié* où la particule fictive évolue en oscillant entre les deux valeurs extrêmes r_2 et r'_2 qui sont appelées distances apsidales.
- $E_0 = E_1$ alors $r = r_1$, on a un mouvement circulaire si $\dot{\theta} \neq 0$ ou un point d'équilibre si $\dot{\theta} = 0$. On remarque que l'équilibre est alors stable : un développement de Taylor au voisinage du point d'équilibre conduit à l'équation d'une parabole

d'équation : $U_{\text{eff}} = K \frac{(r - r_1)^2}{2} + E_1$, avec $K = \left(\frac{d^2 U_{\text{eff}}}{dr^2} \right)_{r=r_1} > 0$ du fait qu'on a un puits de potentiel, ce qui traduit la stabilité de l'équilibre.

- $E_0 < E_1$ alors aucune solution n'est possible et il ne peut donc y avoir de mouvement.

5.3 Barrière d'énergie potentielle effective

Dans ce cas, la courbe donnant l'énergie potentielle effective en fonction de r passe par un maximum.

Sa valeur limite quand $r \rightarrow +\infty$ est également zéro pour les mêmes raisons que précédemment.

On obtient par exemple l'allure suivante :

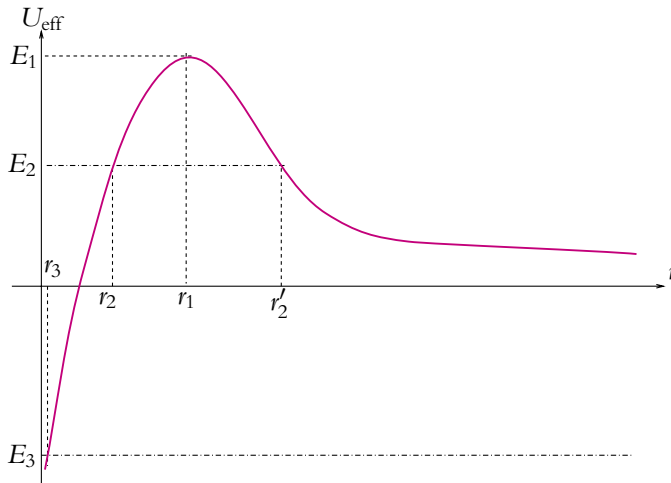


Figure 30.3 Barrière d'énergie potentielle effective.

On a toujours

$$E_0 - U_{\text{eff}}(r) \geq 0$$

On distingue donc plusieurs types de mouvement suivant la valeur de E_0 :

- $E_0 \leq 0$ (par exemple $E_0 = E_3$ sur le schéma ci-dessus) : l'inégalité précédente est vérifiée si $r \leq r_3$ d'après le graphe. On a alors un état lié entre 0 et r_3 .
- $0 \leq E_0 \leq E_1$ (par exemple $E_0 = E_2$ sur le schéma ci-dessus) : l'inégalité précédente est vérifiée si $r \leq r_2$ ou $r \geq r'_2$ d'après le graphe. On a alors un état lié si $r \leq r_2$ entre 0 et r_2 ou un état libre si $r \geq r'_2$.
- $E_0 = E_1$ alors $r = r_1$, on a un point d'équilibre si $\dot{\theta} = 0$. On remarque que l'équilibre est alors instable et que le mouvement circulaire est ici impossible : un

développement de Taylor au voisinage du point d'équilibre conduit à l'équation d'une parabole d'équation : $U_{\text{eff}} = K \frac{(r - r_1)^2}{2} + E_1$, avec $K = \left(\frac{d^2 U_{\text{eff}}}{dr^2} \right)_{r=r_1} < 0$ du fait qu'on a une barrière de potentiel, ce qui traduit l'instabilité de l'équilibre.

- $E_0 > E_1$ alors toutes les valeurs de r sont possibles et on a un état libre.

5.4 Intérêt de cette approche

L'intérêt de ce qui vient d'être fait est de pouvoir discuter la nature du mouvement à partir de la seule expression de l'énergie potentielle effective. On n'a pas forcément besoin de déterminer la trajectoire pour connaître un certain nombre d'informations sur le système.

Cela représente une technique très pratique surtout lorsqu'il n'est pas simple d'obtenir une équation précise de la trajectoire notamment pour des potentiels efficaces plus complexes que ceux qui seront vus lors de l'étude du problème de Kepler et du mouvement des planètes.

A. Application directe du cours

1. Promenade en bateau

Deux personnes de 70 kg sont dans un bateau de 120 kg. Ils descendent une rivière avec une vitesse rectiligne de 2 m.s^{-1} par rapport aux rives.

1. L'une des personnes tombe à l'arrière du bateau. Que devient la vitesse de ce dernier ?
2. Même question si cette personne plonge dans la rivière avec une vitesse de 3 m.s^{-1} par rapport au rivage dans le sens opposé à celui du bateau.
3. Même question si le plongeur a lieu perpendiculairement au mouvement du bateau.

2. Étoiles doubles

On considère deux étoiles formant dans l'espace une "étoile double" que l'on pourra considérer comme isolée. Sous l'action de leur attraction mutuelle, ces deux étoiles décrivent des orbites circulaires autour du centre d'inertie de l'ensemble. On observe que les rayons des orbites sont dans le rapport $\alpha = 2$, que la distance entre les deux étoiles est $L = 3.10^{12} \text{ m}$ et que la période de la révolution de ce système est $T = 50$ années. En déduire la masse de chaque étoile. On donne $G = 6,67.10^{-11} \text{ S.I.}$.

3. Interaction entre deux charges

A et B sont deux particules identiques de charge q et de masse m distantes de D .

1. Les deux charges initialement immobiles sont libérées. Déterminer leur vitesse lorsqu'elles seront infiniment éloignées.
2. A est initialement immobile, B est lancée vers A à la vitesse V_0 . Quelle est la distance minimale entre les deux charges ?

4. Deux points liés par une tige sans masse

Deux points matériels A et B de même masse m sont liés par une tige rigide de masse négligeable, de longueur L . Ces points peuvent se déplacer sans frottement sur un plan horizontal xOy . Initialement, A et B sont sur l'axe Ox , le point A étant en O . La vitesse initiale de A est $\vec{v}_0 = v_0 \vec{u}_y$ perpendiculaire à Ox , alors que celle de B est nulle.

1. Déterminer le mouvement ultérieur de G , centre de gravité des deux points.
2. En utilisant le mobile fictif M dans le référentiel barycentrique $\mathcal{R}_B$, déterminer le mouvement de A et B , dans $\mathcal{R}_B$.

B. Exercices et problèmes

1. Modèle de Bohr de l'atome d'hydrogène

Bohr a proposé le modèle suivant pour décrire la constitution de l'atome d'hydrogène : l'électron est animé d'un mouvement circulaire uniforme autour du noyau de centre d'inertie N

et le moment cinétique en N est quantifié par $\sigma = n \frac{h}{2\pi}$ où n est un entier et h la constante de Planck. On notera qu'il ne s'agit que d'un modèle et que ce dernier n'est pas entièrement satisfaisant : il est nécessaire de recourir à la mécanique quantique pour décrire correctement l'atome.

1. Rappeler l'expression de la force d'interaction exercée par le noyau sur l'électron.
2. Exprimer le carré de la vitesse v de l'électron en fonction de la distance r de ce dernier au noyau.
3. En déduire l'expression du rayon de la trajectoire en fonction de n , h , m_e , e .
4. Calculer sa valeur pour $n = 1$.
5. Exprimer l'énergie totale E de l'électron et montrer qu'elle se met sous la forme $E = -\frac{A}{n^2}$. Donner l'expression et la valeur numérique de A , en électron-volts (on rappelle que $1 \text{ eV} = 1,6 \cdot 10^{-19} \text{ J}$).
6. Sachant que le passage d'un niveau d'énergie à un autre se traduit par l'émission d'un photon de fréquence ν telle que $\Delta E = h\nu$, en déduire que les longueurs d'onde λ émises vérifient : $\frac{1}{\lambda} = R_H \left(\frac{1}{n_2^2} - \frac{1}{n_1^2} \right)$. On rappelle que $\nu = \frac{c}{\lambda}$ où c désigne la vitesse de la lumière dans le vide ($c = 3 \cdot 10^8 \text{ m.s}^{-1}$). On donnera l'expression de la constante de Rydberg R_H ainsi que sa valeur numérique.
7. Quelles sont les longueurs d'onde dans le visible pour les séries de Lyman ($n_2 = 1$) et de Balmer ($n_2 = 2$) ?

2. Etude d'une étoile double, d'après ESTP-ENSAM 1998

On considère dans ce problème une étoile double constituée de deux éléments assimilables à deux points matériels M_1 et M_2 de masse respective m_1 et m_2 . Dans tout le problème, on suppose que le champ de gravitation dû aux autres constituants de la galaxie à laquelle appartient cette étoile double est négligeable.

1. Généralités :

- a) Rappeler l'énoncé précis du principe d'inertie.
- b) Définir le référentiel barycentrique d'un système de points matériels.
- c) Quelle est la propriété vérifiée par le référentiel barycentrique de l'étoile double considérée ici ? Justifier la réponse.

2. Mouvement des deux éléments dans le référentiel barycentrique :

- a) Établir l'expression des quantités de mouvements de M_1 et de M_2 dans le référentiel barycentrique en fonction de la masse réduite μ du système et de la vitesse relative $\vec{v}$.
- b) Écrire l'expression du moment cinétique du système au centre de masse G du système dans le référentiel barycentrique en fonction de $\vec{GM}$, de μ et de $\vec{v}$. On précisera la position et la signification du point M . Quelle est la propriété particulière de ce moment cinétique ?

- c) Établir l'expression de l'énergie cinétique dans le référentiel barycentrique en fonction de μ et $\vec{v}$.
- d) Montrer que le mouvement du point M est celui d'une masse ponctuelle μ soumise à une force égale à l'attraction exercée par M_1 sur M_2 .
- e) Indiquer comment on déduit les trajectoires de M_1 et M_2 à partir de celle de M .

3. Étude du mouvement de M :

- a) Montrer que la trajectoire de M est plane.
- b) Établir qu'elle vérifie la loi des aires qu'on explicitera. On notera C la constante des aires.
- c) On se place alors dans une base telle que la trajectoire de M soit dans le plan d'équation $y = 0$ et on repère le point M par ses coordonnées polaires avec θ l'angle entre le vecteur directeur de Oz et le vecteur position.

Montrer que l'équation du mouvement peut s'écrire :

$$\mu \frac{d\vec{v}}{dt} = \frac{Gm_1m_2}{C} \frac{d\vec{u}_\theta}{dt}$$

en notant G la constante de gravitation universelle.

- d) En déduire l'expression de $\vec{v}$.
- e) Montrer que, moyennant un choix *ad'hoc* des axes Gx et Gz , la trajectoire de M est une conique de foyer G et d'équation $\frac{p}{1 + e \cos \theta}$.
- f) Exprimer p en fonction de G , de C et de $m = m_1 + m_2$.
- g) Dans la base choisie précédemment, exprimer les équations des trajectoires de M_1 et M_2 ainsi que de leurs vitesses $\vec{v}_1$ et $\vec{v}_2$. On prendra $\theta_1 = \theta_2 = \theta$.

4. Observation du phénomène :

Un observateur lié au référentiel barycentrique et très éloigné de G reçoit la lumière émise par chacune des deux constituants de l'étoile double. Le sens et la direction de propagation de la lumière sont définis dans le référentiel barycentrique par un vecteur unitaire $\vec{u}$ parallèle au plan d'équation $z = 0$ et faisant un angle φ avec l'axe Ox . On cherche à observer l'émission de la raie α de l'hydrogène de longueur d'onde $\lambda_0 = 656,3$ nm. Le mouvement des sources d'émission décale cette longueur d'onde par effet Doppler et on observe pour le constituant i une longueur d'onde

$$\lambda_i = \lambda_0 \left(1 - \frac{\vec{v}_i \cdot \vec{u}}{c} \right)$$

en notant c la vitesse de la lumière.

- a) Faire un schéma des vecteurs unitaires et exprimer $\vec{u}$ dans la base $\vec{u}_r, \vec{u}_\theta, \vec{u}_\varphi$.

- b) Exprimer les longueurs d'onde λ_1 et λ_2 correspondant aux deux éléments de l'étoile double en fonction de λ_0 , G , m_1 , m_2 , φ , c , C , e et θ .
- c) Des mesures montrent que ces longueurs d'onde sont des fonctions sinusoïdales du temps. En déduire la nature de la trajectoire de M et montrer qu'elle s'effectue à vitesse angulaire ω constante.
- d) On définit l'amplitude relative des variations de la longueur d'onde λ_i observée par :

$$\epsilon_i = \frac{\lambda_{i,\max} - \lambda_{i,\min}}{2\lambda_0}$$

en notant $\lambda_{i,\max}$ et $\lambda_{i,\min}$ les valeurs maximale et minimale de λ_i . Exprimer ϵ_1 et ϵ_2 en fonction de m_1 , m_2 , C , c , G , ω et φ .

- e) Des mesures montrent que $\epsilon_1 = \epsilon_2$. Que peut-on en déduire sur les masses m_1 et m_2 ?
- f) Dans ces conditions, exprimer les vitesses de M_1 et M_2 dans le référentiel barycentrique en fonction de G , de m , de $\vec{u}_\theta$ et de T la période de révolution du système.
- g) En déduire l'expression de m la masse totale du système en fonction de G , de c , de T , de ϵ et de φ .
- h) Calculer numériquement les valeurs de m , de v_1 et de v_2 . On donne $c = 3.10^8 \text{ m.s}^{-1}$, $G = 6,67.10^{-11} \text{ m}^3.\text{kg}^{-1}.\text{s}^{-2}$, $\epsilon = 9.10^{-4}$, $T = 165$ jours et $\varphi = 0$.
- i) Exprimer la taille caractéristique du système $\rho = M_1 M_2$ en fonction de G , m et T .
- j) Calculer la valeur numérique de ρ .

3. Forces intermoléculaires, d'après Centrale TSI 2003

Dans le cas d'un gaz suffisamment dilué, la force d'interaction $\vec{F}$ entre deux molécules voisines neutres M_1 et M_2 , non polaires, est appelée force de Van der Waals. Lennard-Jones a proposé en 1929 l'expression suivante de l'énergie potentielle pour cette force :

$$E_p(r) = \frac{a}{r^{12}} - \frac{b}{r^6}$$

où r représente la distance entre les deux molécules.

Pour le dioxyde de carbone (CO_2) : $a = \epsilon\sigma^{12}$ et $b = \epsilon\sigma^6$ avec :

- $\epsilon = 9.10^{-2} \text{ eV}$ sachant que $1 \text{ eV} = 1,6.10^{-19} \text{ J}$,
- $\sigma = 3,7.10^{-10} \text{ m}$.

1. Exprimer la force d'interaction $\vec{F}$. Dans quel domaine de distances la force d'interaction est-elle, attractive, répulsive, nulle ? Déterminer les positions d'équilibre.
2. L'étude d'un tel système formé de deux particules isolées s'effectue dans le référentiel barycentrique $\mathcal{B}$.
 - a) Rappeler la définition de ce référentiel.
 - b) Est-il galiléen ?

3. Soit M une particule fictive de masse μ telle que $1/\mu = 1/m + 1/m = 2/m$, où $m = 7,31 \cdot 10^{-26}$ kg désigne la masse d'une molécule de CO_2 . Sa position est définie dans $\mathcal{B}$ par le vecteur $\vec{r} = \vec{GM} = \vec{M_1M_2} = r\vec{u}_r$, G représentant le centre de masse des deux molécules.

a) Exprimer les positions $\vec{r}_1 = \vec{GM_1}$ et $\vec{r}_2 = \vec{GM_2}$ de deux molécules en fonction de $\vec{r}$.

b) Par quelle transformation mathématique passe-t-on de la trajectoire de M à celles de M_1 et M_2 dans le référentiel barycentrique ? Que peut-on en conclure quant aux formes des ces trajectoires ?

4. Montrer que la quantité de mouvement barycentrique de M_2 s'exprime très simplement en fonction de celle de M .

5. a) En appliquant le principe fondamental de la dynamique à M_2 dans $\mathcal{B}$, montrer que tout se passe comme si M était soumise à la force $\vec{F}$ calculée à la question (1).

b) Montrer que le moment cinétique $\vec{L}_b = \vec{GM} \wedge \vec{v}$ dans $\mathcal{B}$ est constant. En déduire que le mouvement est plan.

c) Exprimer la vitesse de M dans le repère en coordonnées polaires (r, θ) de centre G .

6. a) Définir et exprimer en fonction de r , μ , E_p et L_b l'énergie potentielle effective (ou efficace) $E_{p,\text{eff}}$.

b) Montrer dans le cas du potentiel de Lennard-Jones, en effectuant un changement de variable linéaire que $E_{p,\text{eff}}$ peut s'exprimer en unité de ϵ sous la forme :

$$\frac{E_{p,\text{eff}}(x)}{\epsilon} = \left(\frac{1}{x}\right)^{12} - \left(\frac{1}{x}\right)^6 + \frac{p}{x^2}$$

Exprimer la nouvelle variable x en fonction de L_b et r .

c) Exprimer $\dot{r}(x)$.

d) Vérifier que p est une fonction simple de L_b , μ , ϵ et σ .

e) Calculer L_b pour $p = 0,2$ et $p = 0,8$.

La figure (30.4) représente la variation de $E_{p,\text{red}} = E_{p,\text{eff}}/\epsilon$ en fonction de x . La figure (30.5) représente les variations de $y = \sqrt{\mu} \dot{r}$ en fonction de x : c'est ce qu'on appelle un portrait de phase. Sur ce dernier graphique, les courbes en trait fin correspondent à $p = 0,2$ et les courbes en trait épais à $p = 0,8$.

7. On note E l'énergie mécanique du point M . Dans le cas $p = 0,2$, on désigne par E_0 le minimum et $E_3 = 0,0348$ le second extremum (maximum). L'énergie E est en unité de ϵ . On précisera dans les discussions qui suivent si on a affaire à un état de diffusion ou à un état lié.

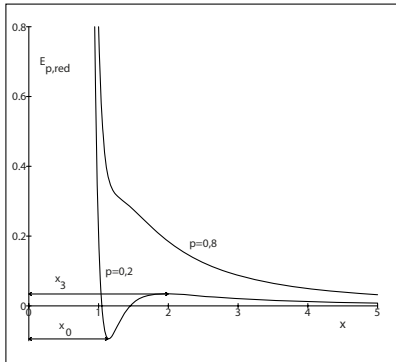


Figure 30.4

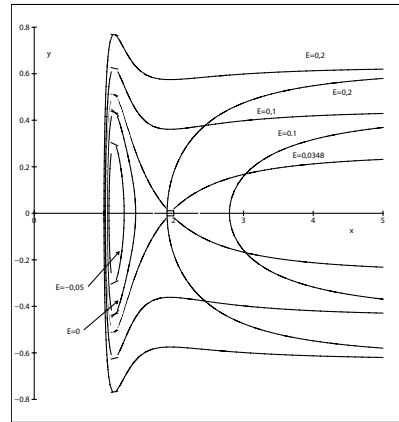


Figure 30.5

- a) Que se passe-t-il pour $E = E_0$? En particulier que vaut $\dot{r}$? Quelle est la nature du mouvement de M (trajectoire, vitesse angulaire, ...)?
- b) On étudie maintenant la nature du mouvement de M pour une valeur E_1 de E telle que $E_0 < E_1 \leq 0$. Quelle est l'état possible pour M ? Préciser sur un schéma les valeurs extrémales de x .
- c) On choisit une valeur E_2 de E telle que $0 < E_2 < E_3$. Quels sont les états possibles de M ?
- d) Montrer que $E = E_3$ correspond à quatre états possibles. Donner l'allure des trajectoires de M .
- e) Étudier de même la situation $E = E_4 > E_3$.
- f) Décrire la situation pour $p = 0, 8$.

Mouvement à force centrale et potentiel newtonien

31

L'objet de ce chapitre est l'étude du mouvement d'un point matériel soumis à une force centrale conservative. On a vu dans le cas du mouvement de deux points matériels isolés en interaction entre eux que l'étude peut se ramener à celle du mouvement d'une particule fictive soumise à la force d'interaction. Dans de nombreux cas, cette force est conservative et dérive d'un potentiel. C'est à ce type d'exemples qu'on s'intéresse ici.

1. Force centrale conservative

1.1 Définitions

On dit qu'un point matériel est soumis à une *force centrale* si cette force est constamment dirigée vers un point fixe O dans le référentiel considéré.

Elle sera dite conservative si son travail peut s'écrire sous la forme de l'opposé d'une différentielle :

$$\delta W = -dEp$$

Comme $\vec{f}$ est constamment dirigée vers un point fixe O , on peut choisir ce point comme origine. La force n'a alors de composante que suivant $\vec{u}_r$: $\vec{f} = f\vec{u}_r$. Le déplacement élémentaire s'écrit : $d\vec{OM} = dr\vec{u}_r + r d\theta \vec{u}_\theta + r \sin \theta d\varphi \vec{u}_\varphi$ en coordonnées sphériques. On en déduit l'expression du travail élémentaire :

$$\delta W = f dr$$

Si la force est conservative,

$$f dr = -dEp$$

On en déduit que l'énergie potentielle E_p n'est fonction que de r , distance du point matériel à l'origine ou point fixe O :

$$\vec{f} = -\frac{dE_p}{dr} \vec{u}_r$$

Exemples

C'est le cas de l'interaction gravitationnelle :

$$E_p = -G \frac{m_1 m_2}{r}$$

en supposant qu'une masse ponctuelle est située en O et l'autre en M , ou de l'interaction coulombienne :

$$E_p = \frac{1}{4\pi\epsilon_0} \frac{q_1 q_2}{r}$$

en supposant qu'une charge ponctuelle est située en O et l'autre en M .

C'est également le cas de l'interaction avec un ressort, de longueur à vide r_0 , fixé en un point fixe O :

$$E_p = \frac{1}{2} k (r - r_0)^2$$

Dans la suite de ce paragraphe, on ne spécifie pas la dépendance de l'énergie potentielle avec r et on établit un certain nombre de résultats généraux des mouvements à force centrale conservative.

1.2 Conservation du moment cinétique

On applique le théorème du moment cinétique en O (qui est un point fixe) :

$$\frac{d\vec{L}_O}{dt} = \vec{OM} \wedge \vec{f} = \vec{0}$$

car le produit vectoriel de deux vecteurs colinéaires est nul, et, par définition des forces centrales, $\vec{f}$ est colinéaire à $\vec{OM}$.

On en déduit que le moment cinétique par rapport au point fixe O est constant.

Le moment cinétique par rapport au point fixe O vers lequel est dirigée la force centrale :

$$\vec{L}_O = \vec{OM} \wedge m \vec{v}$$

se conserve au cours du mouvement.

1.3 Mouvement plan

La conséquence immédiate de la conservation du moment cinétique est la nature plane du mouvement. En effet, le mouvement a lieu perpendiculairement à $\vec{L}_O$ donc dans le plan passant par O et perpendiculaire à $\vec{L}_O$, soit le plan défini par O , la position initiale M_0 et la vitesse initiale $\vec{v}_0$.

1.4 Constante des aires

La deuxième conséquence est connue sous le nom de loi des aires. Du fait de sa nature plane, le mouvement va être décrit en coordonnées polaires dans le plan du mouvement. On a donc :

$$\begin{aligned}\vec{OM} &= r\vec{u}_r \\ \vec{v} &= \dot{r}\vec{u}_r + r\dot{\theta}\vec{u}_\theta\end{aligned}$$

On en déduit l'expression du moment cinétique :

$$\vec{L}_O = \vec{OM} \wedge m\vec{v} = mr^2\dot{\theta}\vec{u}_z$$

Ce vecteur étant constant d'après ce qui précède, on en déduit que

$$mr^2\dot{\theta} = \text{constante}$$

On introduit la constante des aires qu'on note C :

$$C = r^2\dot{\theta}$$

qui correspond au moment cinétique par rapport au point central O par unité de masse. Il s'agit d'une constante, car la masse d'un point matériel ne varie pas.

Cette constante peut s'interpréter en termes de vitesse aréolaire. On appelle vitesse aréolaire $\mathcal{V}$ la vitesse à laquelle le rayon vecteur balaie l'aire définie par la trajectoire dans le plan du mouvement.

On assimile le triangle OM_1M_2 à un triangle rectangle en M_1 :

$$\mathcal{V} = \frac{dA}{dt} = \frac{r(rd\theta)}{2dt} = \frac{r^2\dot{\theta}}{2} = \frac{C}{2}$$

qui est donc une constante.

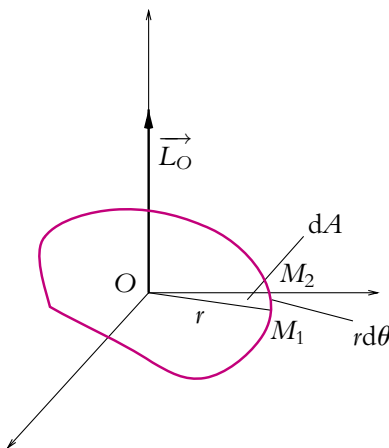


Figure 31.1 Constante des aires et vitesse aréolaire.

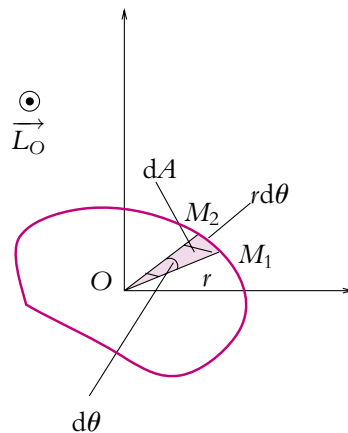


Figure 31.2 Projection dans le plan du mouvement.

1.5 Conservation de l'énergie mécanique

Si le système ne subit aucune autre force que la force centrale conservative, il y a conservation de l'énergie mécanique.

On peut le traduire en coordonnées polaires puisque ce sont les coordonnées adaptées au type de mouvement considéré.

$$Em = Ec + Ep$$

Or on a déjà établi que l'énergie potentielle ne dépend que de r . D'autre part, l'énergie cinétique s'écrit :

$$Ec = \frac{1}{2}mv^2 = \frac{1}{2}m(\dot{r}^2 + r^2\dot{\theta}^2)$$

En introduisant la constante des aires C :

$$\dot{\theta} = \frac{C}{r^2}$$

on obtient :

$$Ec = \frac{1}{2}m\dot{r}^2 + \frac{1}{2}mr^2\frac{C^2}{r^4} = \frac{1}{2}m\dot{r}^2 + \frac{1}{2}m\frac{C^2}{r^2}$$

et l'énergie mécanique s'écrit :

$$Em = \frac{1}{2}m\dot{r}^2 + \frac{mC^2}{2r^2} + Ep(r) = \text{constante}$$

1.6 Énergie potentielle effective et étude qualitative du mouvement

On introduit donc une énergie potentielle effective :

$$Ep_{\text{eff}}(r) = \frac{mC^2}{2r^2} + Ep(r)$$

et on a :

$$\frac{1}{2}m\dot{r}^2 + Ep_{\text{eff}}(r) = \text{constante}$$

Comme cela a été fait au paragraphe 5 du chapitre 30 sur le mouvement de deux points matériels isolés, il est possible de discuter qualitativement la nature du mouvement à partir de l'allure de l'énergie potentielle effective en fonction de r et de la valeur de l'énergie mécanique. En effet, $\frac{1}{2}m\dot{r}^2 > 0$ donc

$$Em > Ep_{\text{eff}}(r)$$

ce qui restreint les possibilités de mouvement suivant la position de Em par rapport à la courbe représentative de l'énergie potentielle effective.

On distingue :

- les états liés pour lesquels le point matériel évolue entre deux positions $r_{\min}$ et $r_{\max}$,
- les états libres pour lesquels le point matériel peut aller à l'infini.

1.7 Trajectoire

De la conservation de l'énergie mécanique, on peut déduire la relation entre θ et r , c'est-à-dire l'équation de la trajectoire en coordonnées polaires :

$$\begin{cases} Em = \frac{1}{2}m\dot{r}^2 + Ep_{\text{eff}}(r) \\ C = r^2\dot{\theta} \end{cases}$$

alors

$$\frac{d\theta}{dr} = \frac{\dot{\theta}}{\dot{r}} = \frac{\frac{C}{r^2}}{\pm \sqrt{\frac{2(Em - Ep_{\text{eff}}(r))}{m}}} = \pm \frac{C}{r^2} \sqrt{\frac{m}{2(Em - Ep_{\text{eff}}(r))}}$$

On sépare les deux cas suivant le signe de $\frac{d\theta}{dr}$. Connaissant alors l'expression de $Ep_{\text{eff}}(r)$, on peut « théoriquement » obtenir :

$$\theta - \theta_0 = \int_{r_0}^r \frac{C}{r^2} \sqrt{\frac{m}{2(Em - Ep_{\text{eff}}(r))}} dr \quad \text{quand} \quad \frac{d\theta}{dr} > 0$$

ou

$$\theta - \theta_0 = - \int_{r_0}^r \frac{C}{r^2} \sqrt{\frac{m}{2(Em - Ep_{\text{eff}}(r))}} dr \quad \text{quand} \quad \frac{d\theta}{dr} < 0$$

où $\theta_0 = \theta(r_0)$.

Cela ne dit pas que le calcul donnera une expression analytique ni qu'il sera simple, mais il est *a priori* possible d'obtenir l'expression de la trajectoire en coordonnées polaires.

On peut également procéder de même pour obtenir une expression de r en fonction du temps : $\dot{r} = \frac{dr}{dt}$. Donc

$$t - t_0 = \pm \int_{r_0}^r \frac{dr}{\sqrt{\frac{2(Em - Ep_{\text{eff}}(r))}{m}}}$$

où $r_0 = r(t_0)$. On inverse la relation obtenue pour avoir $r(t)$. Alors on peut déterminer l'équation horaire $\theta(t)$ par :

$$\theta - \theta_0 = \int_{t_0}^t \frac{C}{r^2} dt$$

2. Potentiel newtonien

2.1 Interaction newtonienne

Une interaction est dite *newtonienne* si la force d'interaction associée varie en fonction de r selon une loi en $\frac{1}{r^2}$:

$$\vec{f} = \frac{K}{r^2} \vec{u}_r$$

Ainsi les interactions gravitationnelles :

$$\vec{f} = -\frac{Gm_1m_2}{r^2} \vec{u}_r$$

ou coulombiennes :

$$\vec{f} = \frac{q_1q_2}{4\pi\epsilon_0r^2} \vec{u}_r$$

sont des interactions newtoniennes.

2.2 Expression de l'énergie potentielle newtonienne

La force d'interaction newtonienne dérive d'un potentiel du type :

$$Ep = \frac{K}{r} + \text{cste}$$

En effet, le travail élémentaire de cette force s'écrit :

$$\delta W = \vec{f} \cdot d\vec{r} = \frac{K}{r^2} dr = -d\left(\frac{K}{r}\right)$$

soit une énergie potentielle correspondante :

$$\delta W = -dEp$$

avec :

$$Ep = \frac{K}{r} + \text{cste}$$

Dans le cas des interactions gravitationnelles, on a :

$$Ep = -\frac{Gm_1m_2}{r} + \text{cste}$$

et pour les interactions coulombiennes :

$$Ep = \frac{q_1q_2}{4\pi\epsilon_0r} + \text{cste}$$

On choisit l'origine des potentiels nulle à l'infini, ce qui traduit le fait que des points matériels infiniment éloignés n'exercent pas d'interaction entre eux. La constante qui

apparaît dans les expressions de l'énergie potentielle est donc nulle. On prendra donc :

$$E_p = \frac{K}{r}$$

2.3 Interaction attractive ou répulsive

L'interaction sera attractive si $K < 0$ et répulsive si $K > 0$.

L'interaction gravitationnelle est donc toujours attractive, c'est pourquoi on parle d'attraction gravitationnelle. En revanche, l'interaction coulombienne est attractive pour deux charges de signe opposé et répulsive pour deux charges de même signe.

2.4 Discussion qualitative du mouvement

On déduit donc l'expression de l'énergie potentielle effective dans le cas d'une interaction newtonienne :

$$E_{p_{\text{eff}}}(r) = \frac{mC^2}{2r^2} + E_p(r) = \frac{mC^2}{2r^2} + \frac{K}{r}$$

a) Interaction attractive

Dans ce cas, $K < 0$.

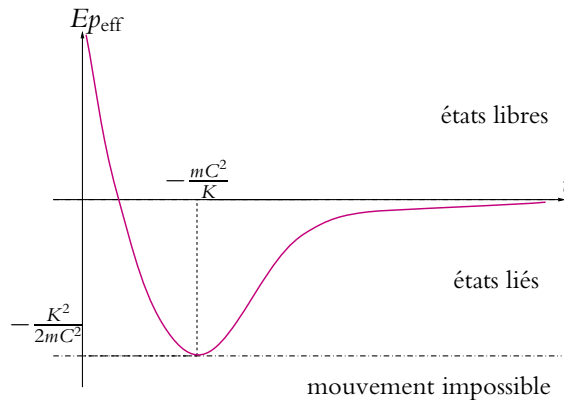


Figure 31.3 Énergie potentielle effective pour une interaction attractive.

On étudie tout d'abord l'énergie potentielle effective comme fonction de r :

$$f(r) = \frac{mC^2}{2r^2} + \frac{K}{r} \quad \text{donc} \quad f'(r) = -\frac{mC^2}{r^3} - \frac{K}{r^2} = -\frac{mC^2 + Kr}{r^3}$$

$f'(r)$ s'annule pour $r = -\frac{mC^2}{K} > 0$ puisque $K < 0$. $f'(r)$ est positif pour $r > -\frac{mC^2}{K}$.

On a donc un minimum en $r = -\frac{mC^2}{K}$.

On obtient l'allure de la figure 31.3 puisque :

$$\lim_{r \rightarrow +\infty} Ep_{\text{eff}}(r) = 0 \quad \text{et} \quad \lim_{r \rightarrow 0} Ep_{\text{eff}}(r) = +\infty$$

On en déduit l'existence d'états libres et d'états liés dans le cas d'un potentiel newtonien attractif.

b) Interaction répulsive

Dans ce cas, $K > 0$.

$f'(r)$ ne s'annule pas puisque $K > 0$ et reste constamment négatif. On a donc une fonction décroissante.

On obtient l'allure de la figure 31.4 puisque :

$$\lim_{r \rightarrow +\infty} Ep_{\text{eff}}(r) = 0$$

et

$$\lim_{r \rightarrow 0} Ep_{\text{eff}}(r) = +\infty$$

On en déduit qu'il n'y a que des états libres possibles dans le cas d'un potentiel newtonien répulsif.

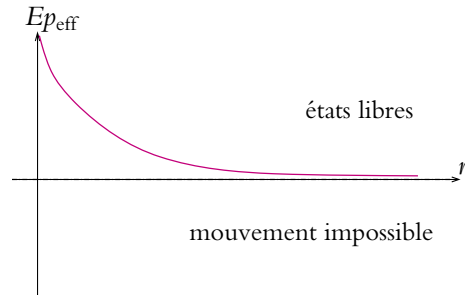


Figure 31.4 Énergie potentielle effective pour une interaction répulsive.

3. Détermination de l'expression de la trajectoire (MPSI, PCSI)

3.1 Établissement de la trajectoire à partir du vecteur excentricité

a) Vecteur excentricité et loi de conservation

Le principe fondamental de la dynamique s'écrit ici :

$$m \frac{d\vec{v}}{dt} = \frac{K}{r^2} \vec{u}_r$$

Or on sait que $\vec{u}_r = -\frac{d\vec{u}_\theta}{d\theta} = -\frac{1}{\dot{\theta}} \frac{d\vec{u}_\theta}{dt}$, soit

$$\frac{d\vec{v}}{dt} = -\frac{K}{mr^2 \dot{\theta}} \frac{d\vec{u}_\theta}{dt} = -\frac{K}{mC} \frac{d\vec{u}_\theta}{dt}$$

et par intégration :

$$\vec{v} = -\frac{K}{mC} \vec{u}_\theta + \vec{c}^{\text{te}} = -\frac{K}{mC} (\vec{u}_\theta + \vec{e})$$

où $\vec{c}^{\text{te}}$ et $\vec{e}$ sont des vecteurs constants.

On définit à partir de cette équation la conservation du vecteur excentricité :

$$\vec{e} = -\vec{u}_\theta - \frac{mC}{K} \vec{v}$$

Il s'agit d'un vecteur dont la norme est sans dimension qui est conservé dans le type de mouvement étudié ici.

b) Établissement de la trajectoire

L'intérêt de ce vecteur est de permettre de déterminer assez simplement l'équation de la trajectoire.

On multiplie scalairement le vecteur excentricité, qui est constant d'après le paragraphe précédent, par $\vec{u}_\theta$:

$$-\vec{e} \cdot \vec{u}_\theta = \left(\vec{u}_\theta \cdot \vec{u}_\theta + \frac{mC}{K} \vec{v} \cdot \vec{u}_\theta \right) = 1 + \frac{mC}{K} v_\theta$$

où v_θ désigne la composante orthoradiale de la vitesse.

$$v_\theta = \frac{K}{mC} (-\vec{e} \cdot \vec{u}_\theta - 1) = r\dot{\theta} = \frac{C}{r}$$

soit

$$\frac{1}{r} = \frac{K}{mC^2} (-\vec{e} \cdot \vec{u}_\theta - 1)$$

Le vecteur $\vec{e}$ étant constant, son module e et sa direction le sont aussi. On note θ_0 l'angle entre l'axe Oy et $\vec{e}$. L'angle entre les vecteurs $\vec{e}$ et $\vec{u}_\theta$ est $\theta - \theta_0$ donc $\vec{e} \cdot \vec{u}_\theta = e \cos(\theta - \theta_0)$.

L'équation de la trajectoire est donc :

$$r = \frac{\frac{-mC^2}{K}}{1 + e \cos(\theta - \theta_0)}$$

On peut choisir l'axe Oy colinéaire à la direction du vecteur excentricité puisque ce dernier est constant au cours du mouvement.

Cela revient à choisir $\theta_0 = 0$, ce qui simplifie l'équation de la trajectoire en

$$r = \frac{\frac{-mC^2}{K}}{1 + e \cos \theta}$$

C'est l'équation d'une conique d'excentricité e , ce qui explique le nom donné à ce vecteur. L'origine est l'un des foyers de la conique et son paramètre est $p = \frac{mC^2}{|K|}$.

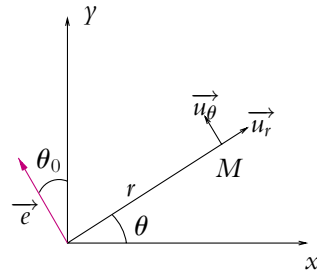


Figure 31.5 Vecteur excentricité.

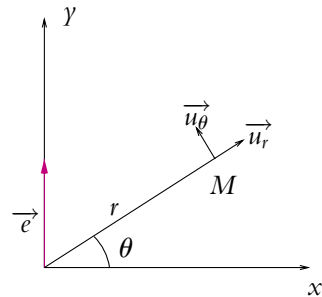


Figure 31.6 Vecteur excentricité et trajectoire.

3.2 Établissement de la trajectoire par la conservation de l'énergie mécanique

Une autre méthode consiste à exploiter la conservation de l'énergie mécanique :

$$Em = Ec + Ep = \frac{1}{2}m\dot{r}^2 + \frac{mC^2}{2r^2} + \frac{K}{r}$$

On cherche l'expression de $r(\theta)$. Sachant que :

$$\dot{r} = \frac{dr}{dt} = \frac{dr}{d\theta} \frac{d\theta}{dt} \quad \text{et} \quad \frac{d\theta}{dt} = \dot{\theta} = \frac{C}{r^2}$$

On déduit

$$Em = \frac{mC^2}{2r^4} \left(\frac{dr}{d\theta} \right)^2 + \frac{mC^2}{2r^2} + \frac{K}{r}$$

Pour des raisons de commodité dans le calcul, on introduit la variable $u(\theta) = \frac{1}{r(\theta)}$, alors :

$$r = \frac{1}{u} \quad \text{donc} \quad \frac{dr}{d\theta} = -\frac{1}{u^2} \frac{du}{d\theta}$$

et

$$Em = \frac{mC^2}{2} u^4 \frac{1}{u^4} \left(\frac{du}{d\theta} \right)^2 + \frac{mC^2}{2} u^2 + Ku = \frac{mC^2}{2} \left(\left(\frac{du}{d\theta} \right)^2 + u^2 \right) + Ku$$

Or Em est une constante donc : $\frac{dEm}{d\theta} = 0$ soit

$$mC^2 \left(\frac{du}{d\theta} \frac{d^2u}{d\theta^2} + u \frac{du}{d\theta} \right) + K \frac{du}{d\theta} = 0$$

Cela donne deux possibilités :

- $\frac{du}{d\theta} = 0$

ce qui impliquerait que u soit constant et que la trajectoire soit un cercle. On verra que cette solution est incluse dans l'autre cas.

- $mC^2 \left(\frac{d^2u}{d\theta^2} + u \right) + K = 0$

soit

$$\frac{d^2u}{d\theta^2} + u = -\frac{K}{mC^2}$$

dont la solution est la somme d'une solution particulière constante et de la solution générale de l'équation homogène associée :

$$u = -\frac{K}{mC^2} + A \cos(\theta - \theta_0)$$

On en déduit :

$$\begin{aligned} \frac{1}{r} &= -\frac{K}{mC^2} + A \cos(\theta - \theta_0) = \frac{-K + mC^2 A \cos(\theta - \theta_0)}{mC^2} \\ \text{soit} \quad r &= \frac{mC^2}{-K + mC^2 A \cos(\theta - \theta_0)} = \frac{\frac{mC^2}{|K|}}{-\text{sign}(K) + A \frac{mC^2}{|K|} \cos(\theta - \theta_0)} \end{aligned}$$

On obtient une équation du type :

$$r = \frac{p}{-\epsilon + e \cos(\theta - \theta_0)}$$

avec $\epsilon = \pm 1$ à savoir l'équation d'une conique de paramètre p et d'excentricité e , O étant l'un des foyers.

3.3 Établissement de la trajectoire à partir de l'invariant de Runge-Lenz ou de Laplace

a) Invariant de Runge-Lenz ou de Laplace

On définit le vecteur de Runge-Lenz ou de Laplace comme la quantité :

$$\vec{R} = \vec{v} \wedge \vec{L}_O + K \vec{u}_r$$

Il s'agit d'un invariant du mouvement étudié ici. En effet :

$$\begin{aligned} \frac{d\vec{R}}{dt} &= \frac{d\vec{v}}{dt} \wedge \vec{L}_O + \vec{v} \wedge \frac{d\vec{L}_O}{dt} + K \frac{d\vec{u}_r}{dt} \\ &= \frac{\vec{f}}{m} \wedge \vec{L}_O + K \dot{\theta} \vec{u}_\theta \\ &= \frac{1}{m} \frac{K}{r^2} \vec{u}_r \wedge m r^2 \dot{\theta} \vec{u}_z + K \dot{\theta} \vec{u}_\theta \\ &= -K \dot{\theta} \vec{u}_\theta + K \dot{\theta} \vec{u}_\theta \end{aligned}$$

soit :

$$\frac{d\vec{R}}{dt} = \vec{0}$$

$\vec{R}$ est donc constant. On peut montrer que

$$\vec{R} = -\frac{K}{mC} \vec{e} \wedge \vec{L}_O$$

où $\vec{e}$ est le vecteur excentricité introduit au paragraphe 3.1.

b) Établissement de la trajectoire

On peut alors retrouver l'équation de la trajectoire à partir de $\vec{R}$:

$$\begin{aligned}\vec{R} \cdot \vec{u}_r &= (\vec{v} \wedge \vec{L}_O) \cdot \vec{u}_r + K \\ &= (\vec{u}_r \wedge \vec{v}) \cdot \vec{L}_O + K \\ &= r \dot{\theta} \vec{u}_z \cdot m C \vec{u}_z + K \\ &= \frac{m C^2}{r} + K\end{aligned}$$

$$\text{Or } \vec{R} \cdot \vec{u}_r = R \cos \varphi$$

en notant φ l'angle entre $\vec{R}$ et $\vec{u}_r$.

$\vec{R}$ étant constant, on peut choisir l'axe Ox selon la direction de $\vec{R}$ alors $\varphi = \theta$ et on

$$\text{retrouve l'équation de la trajectoire : } r = \frac{\frac{m C^2}{K}}{-1 + \frac{R}{K} \cos \varphi}.$$

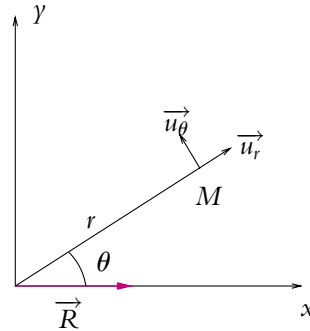


Figure 31.7 Invariant de Runge-Lenz et trajectoire.

3.4 Établissement de la trajectoire par les formules de Binet

a) Formules de Binet

Il s'agit d'exprimer le carré du module de la vitesse et l'accélération en coordonnées polaires en utilisant la variable $u = \frac{1}{r}$.

La vitesse est donnée par : $\vec{v} = \dot{r} \vec{u}_r + r \dot{\theta} \vec{u}_\theta$ soit $v^2 = \dot{r}^2 + r^2 \dot{\theta}^2$

Le fait de poser $u = \frac{1}{r}$ conduit à écrire : $\frac{dr}{d\theta} = -\frac{1}{u^2} \frac{du}{d\theta}$ soit

$$\vec{v} = \frac{dr}{d\theta} \frac{d\theta}{dt} \vec{u}_r + r \frac{d\theta}{dt} \vec{u}_\theta = -\frac{du}{d\theta} \frac{1}{u^2} \dot{\theta} \vec{u}_r + \frac{1}{u} \dot{\theta} \vec{u}_\theta$$

et

$$v^2 = \frac{\dot{\theta}^2}{u^4} \left(\left(\frac{du}{d\theta} \right)^2 + u^2 \right)$$

Or on a un mouvement à force centrale, donc $C = r^2 \dot{\theta} = \frac{\dot{\theta}}{u^2}$ est une constante.

On en déduit :

$$v^2 = C^2 \left(\left(\frac{du}{d\theta} \right)^2 + u^2 \right)$$

qui est la première formule de Binet relative au carré du module de la vitesse. On remarquera que cette formule a déjà été utilisée dans la méthode basée sur la conservation de l'énergie mécanique.

Quant à l'accélération en coordonnées polaires, elle s'écrit :

$$\vec{a} = (\ddot{r} - r\dot{\theta}^2) \vec{u}_r$$

le terme sur $\vec{u}_\theta$ étant nul car la force est centrale.

Avec $u = \frac{1}{r}$:

$$\ddot{r} = \frac{d^2 r}{dt^2} = \frac{d}{dt} \left(\frac{d}{dt} \left(\frac{1}{u} \right) \right) = \frac{d}{dt} \left(\frac{d}{d\theta} \left(\frac{1}{u} \right) \dot{\theta} \right) = \frac{d}{dt} \left(-\frac{1}{u^2} \frac{du}{d\theta} \dot{\theta} \right)$$

Or pour les mouvements à force centrale, on a : $C = r^2 \dot{\theta} = \frac{\dot{\theta}}{u^2}$ constant donc

$$\ddot{r} = -C \frac{d}{dt} \left(\frac{du}{d\theta} \right) = -C \frac{d}{d\theta} \left(\frac{du}{d\theta} \right) \dot{\theta} = -C \dot{\theta} \frac{d^2 u}{d\theta^2} = -C^2 \frac{d^2 u}{d\theta^2}$$

On obtient donc :

$$\vec{a} = \left(-C^2 u^2 \frac{d^2 u}{d\theta^2} - \frac{1}{u} (Cu^2)^2 \right) \vec{u}_r$$

soit

$$\vec{a} = -C^2 u^2 \left(\frac{d^2 u}{d\theta^2} + u \right) \vec{u}_r.$$

Il s'agit de la deuxième formule de Binet relative à l'accélération.

► **Remarque :** On notera que les formules de Binet ne sont valables que pour un mouvement à force centrale et qu'il ne faut pas chercher à l'appliquer dans d'autres circonstances. En effet, l'hypothèse d'une force centrale a été utilisée lors de la démonstration.

b) Obtention de l'équation de la trajectoire

L'application du principe fondamental de la dynamique au point matériel soumis à une force centrale donne :

$$m \vec{a} = \vec{f}$$

$\vec{a}$ et $\vec{f}$ n'ont de composante non nulle que sur $\vec{u}_r$. On projette donc sur cette direction :

$$-mC^2 u^2 \left(\frac{d^2 u}{d\theta^2} + u \right) = \frac{K}{r^2} = Ku^2$$

soit

$$\frac{d^2 u}{d\theta^2} + u = -\frac{K}{mC^2}$$

qui est la même équation différentielle qu'au paragraphe 3.2. On en déduit la solution :

$$u = -\frac{K}{mC^2} + A \cos(\theta - \theta_0) \quad \text{ou} \quad r = \frac{p}{-\epsilon + e \cos(\theta - \theta_0)}$$

3.5 Conclusion

Aucune méthode n'est imposée par le programme : on pourra donc adopter celle que l'on veut. Cependant les trois premières sont préférables car elles reposent sur une loi de conservation du mouvement.

4. Discussion sur la nature du mouvement (MPSI, PCSI)

4.1 À partir de la courbe de l'énergie potentielle effective

Lors de l'analyse qualitative du mouvement, on a distingué :

- les états liés, qui n'existent qu'avec une interaction attractive, et qui correspondent au cas où r ne tend pas vers l'infini. On a vu que cela correspondait à une énergie mécanique négative. On aura alors une ellipse.
- les états libres, qui existent aussi bien avec une interaction attractive que répulsive. Ils correspondent aux cas où r peut tendre vers l'infini. On a vu que cela correspondait à la seule possibilité pour les interactions répulsives ; la trajectoire est la branche d'hyperbole
 - n'enveloppant pas le foyer O pour une interaction répulsive,
 - enveloppant le foyer O pour une interaction attractive.

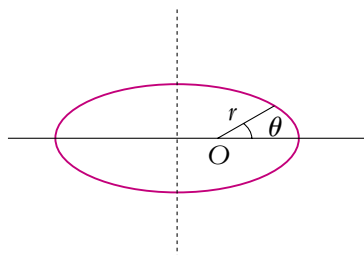


Figure 31.8 Mouvement elliptique.

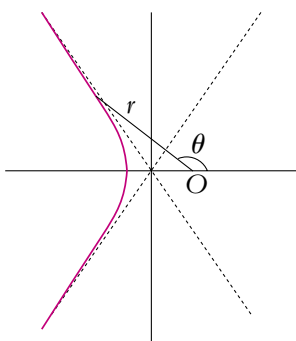


Figure 31.9 Mouvement hyperbolique répulsif.

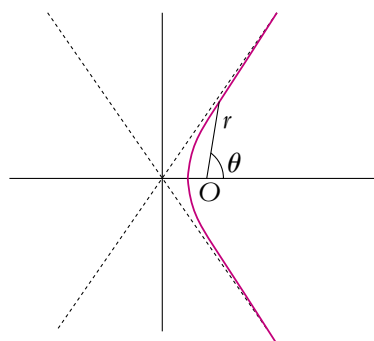


Figure 31.10 Mouvement hyperbolique attractif.

- On aura une parabole pour le cas limite entre ellipse et hyperbole, c'est-à-dire une énergie mécanique nulle pour les interactions attractives.

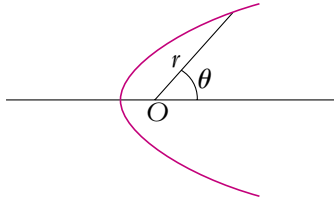


Figure 31.11 Mouvement parabolique.

4.2 À partir de la valeur de l'excentricité

On part de l'équation de la trajectoire obtenue au paragraphe précédent :

$$r = \frac{p}{-\epsilon + e \cos(\theta - \theta_0)}$$

avec

$$p = \frac{mC^2}{|K|} \quad \text{et} \quad \epsilon = \text{sign}(K) = \pm 1$$

D'après les résultats connus sur les coniques (Cf. cours de mathématiques), la trajectoire est :

- une ellipse si $e < 1$ et on aura alors un état lié,
- une hyperbole si $e > 1$ et on aura un état libre,
- une parabole si $e = 1$ et on aura un état libre.

4.3 À partir de la valeur de l'énergie mécanique

On cherche à établir l'expression de l'énergie mécanique en fonction de l'excentricité en calculant le module au carré du vecteur excentricité $\vec{e} = -\left(\vec{u}_\theta + \frac{mC}{K}\vec{v}\right)$:

$$e^2 = \left(\vec{u}_\theta + \frac{mC}{K}\vec{v}\right)^2 = 1 + \frac{m^2 C^2 v^2}{K^2} + 2\frac{mC}{K}\vec{u}_\theta \cdot \vec{v}$$

Or $\vec{u}_\theta \cdot \vec{v} = r\dot{\theta} = \frac{C}{r}$ donc

$$e^2 = 1 + \frac{m^2 C^2 v^2}{K^2} + 2\frac{mC^2}{Kr} = 1 + \frac{2mC^2}{K^2} \left(\frac{1}{2}mv^2 + \frac{K}{r}\right) = 1 + \frac{2mC^2}{K^2} Em$$

On en déduit l'expression de l'énergie mécanique :

$$Em = \frac{K^2}{2mC^2} (e^2 - 1) \quad (31.1)$$

On en déduit la nature de la trajectoire en fonction de la valeur de Em :

- une ellipse si $Em < 0$ (ou $e < 1$),
- une parabole si $Em = 0$ (ou $e = 1$),
- une hyperbole si $Em > 0$ (ou $e > 1$).

ce qui est cohérent avec les paragraphes précédents.

► **Remarque** : On peut obtenir (31.1) à partir du module du vecteur de Runge-Lenz (ou de Laplace) ou encore à partir des formules de Binet.

5. Mouvement des planètes - Lois de Kepler

Dans toute cette partie, on se place dans le cadre de la force de gravitation :

$$\vec{f} = -G \frac{m_1 m_2}{r^2} \vec{u}_r$$

et

$$K = -Gm_1 m_2 < 0$$

On a donc une force attractive et la possibilité d'observer des états libres et des états liés.

5.1 Énoncé des lois de Kepler

Les lois de Kepler ont été énoncées expérimentalement à partir de l'observation du mouvement des planètes. Elles sont au nombre de trois :

1. Le centre des planètes décrit une ellipse dont l'un des foyers est le Soleil.
2. Les rayons vecteurs balaient des aires égales pour des intervalles de temps égaux.
3. Le rapport entre le carré de la période T de révolution de la planète autour du Soleil et le cube du demi-grand axe a de la trajectoire est indépendant de la planète :

$$\frac{a^3}{T^2} = \text{cste}$$

Ces propriétés s'interprètent aisément en postulant les hypothèses suivantes :

- Les planètes et le Soleil présentent une symétrie sphérique permettant de les assimiler à des points matériels. En première approximation, cette hypothèse est justifiée même si l'on sait que les planètes ne sont pas rigoureusement des sphères. Cette approximation pourra éventuellement expliquer des écarts aux lois de Kepler.
- Le mouvement d'une planète est uniquement lié à l'interaction entre cette planète et le Soleil, on exclut toute influence des autres planètes ou objets célestes. Ainsi on se ramène au problème à deux corps. Là encore l'hypothèse est critiquable du fait

même de l'existence des autres planètes et des autres étoiles. On se ramène donc à un mouvement à force centrale pour le mouvement relatif.

- La masse de la planète est négligeable devant celle du Soleil.

L'intérêt de cette dernière hypothèse est de pouvoir assimiler le barycentre du système constitué du Soleil et d'une planète au centre du Soleil. En effet, le barycentre du système est défini par :

$$\overrightarrow{OG} = \frac{M_S \overrightarrow{OS} + M_P \overrightarrow{OP}}{M_S + M_P}$$

en notant M_S et M_P les masses respectives du Soleil et de la planète et S et P leurs positions respectives. Si $M_P \ll M_S$, alors

$$\overrightarrow{OG} \simeq \overrightarrow{OS} + \frac{M_P}{M_S} \overrightarrow{OP} \simeq \overrightarrow{OS}$$

La particule fictive est donc confondue avec la planète. Si ce n'est pas le cas, on travaille avec la particule fictive : M_P est remplacée par $\frac{M_P M_S}{M_P + M_S}$ et $r = GP$ avec G le centre de masse du système formé par le Soleil et la planète.

Les valeurs fournies dans le tableau suivant montrent que cette hypothèse est justifiée du fait des ordres de grandeur sachant que la masse du Soleil vaut $M_S = 2.10^{30}$ kg.

Planètes	Masse (kg)	Pourcentage de la masse solaire
Mercure	$3,2.10^{23}$	$1,6.10^{-5}$
Vénus	$4,9.10^{24}$	$2,5.10^{-4}$
Terre	$6,0.10^{24}$	$3,0.10^{-4}$
Mars	$6,5.10^{23}$	$3,2.10^{-5}$
Jupiter	$1,9.10^{27}$	$9,6.10^{-2}$
Saturne	$5,7.10^{26}$	$2,9.10^{-2}$
Uranus	$8,7.10^{25}$	$4,4.10^{-3}$
Neptune	$1,0.10^{26}$	$5,2.10^{-3}$
Pluton	$1,8.10^{22}$	$9,0.10^{-7}$

On peut noter tout de suite que la seconde loi de Kepler est une conséquence immédiate de la conservation du moment cinétique dans le cadre du système isolé à deux corps envisagé ici. Le mouvement est plan et il vérifie la loi des aires, qui n'est autre que la seconde loi de Kepler.

Comme on se trouve dans le cas d'une interaction attractive, il est possible d'obtenir un état lié. La trajectoire de la planète autour du Soleil sera alors une ellipse d'après

les résultats du paragraphe précédent et l'énergie mécanique sera négative. On va maintenant procéder à une étude approfondie de l'état lié et on pourra alors établir la troisième loi de Kepler dans le cadre des hypothèses qui viennent d'être formulées.

5.2 Étude de la trajectoire circulaire

a) Caractère uniforme du mouvement

Pour une trajectoire circulaire, la distance r au centre du Soleil est constante : $r = R$. L'accélération s'écrit donc :

$$\vec{a} = -R\dot{\theta}^2\vec{u}_r + R\ddot{\theta}\vec{u}_\theta$$

Le principe fondamental de la dynamique s'écrit alors :

$$M_P\vec{a} = M_P(-R\dot{\theta}^2\vec{u}_r + R\ddot{\theta}\vec{u}_\theta) = -G\frac{M_P M_S}{r^2}\vec{u}_r$$

On en déduit les relations suivantes :

$$\begin{cases} -R\dot{\theta}^2 = -\frac{GM_S}{R^2} \\ R\ddot{\theta} = 0 \end{cases}$$

La vitesse angulaire $\omega = \dot{\theta}$ est constante et le mouvement circulaire est uniforme.

b) Relation entre le rayon et la vitesse

La vitesse s'exprimant en fonction de la vitesse angulaire ω par la relation : $v = R\omega$, on en déduit :

$$R\frac{v^2}{R^2} = \frac{GM_S}{R^2}$$

soit

$$R = \frac{GM_S}{v^2} \quad \text{ou} \quad v = \sqrt{\frac{GM_S}{R}}$$

c) Troisième loi de Kepler

Or $v = R\omega = \frac{2\pi R}{T}$ en notant T la période, on déduit de la relation obtenue au paragraphe précédent :

$$R = \frac{GM_S T^2}{4\pi^2 R^2}$$

On retrouve ainsi la troisième de Kepler :

$$\frac{T^2}{R^3} = \frac{4\pi^2}{GM_S}$$

d) Expression des différentes énergies

L'énergie cinétique vaut alors :

$$E_c = \frac{1}{2} M_p v^2 = \frac{1}{2} \frac{G M_p M_S}{R}$$

L'énergie potentielle s'écrit :

$$E_p = -\frac{G M_p M_S}{R} = -2E_c$$

Quant à l'énergie mécanique, elle vaut :

$$E_m = E_c + E_p = -\frac{E_p}{2} + E_p = -\frac{E_p}{2} = -\frac{1}{2} \frac{G M_p M_S}{R}$$

On en déduit que :

$$E_c = -E_m = -\frac{1}{2} E_p$$

5.3 Caractéristiques du mouvement elliptique (MPSI, PCSI)

On a un mouvement elliptique si :

- la force est attractive, ce qui est le cas ici,
- l'énergie mécanique est négative et supérieure à la valeur du minimum de l'énergie potentielle effective. (Cf. figure 31.3).

$$E_m = \frac{1}{2} \mu v_0^2 - \frac{K}{r_0} < 0$$

On choisit ici d'étudier le mouvement du mobile fictif défini au chapitre 22, la masse du système est donc la masse réduite μ .

L'équation est alors :

$$r = \frac{p}{1 + e \cos(\theta - \theta_0)}$$

avec

$$p = \frac{\mu C^2}{|K|}$$

et $e < 1$. On a : $K = -G M_S M_p$ et $\mu = \frac{M_S M_p}{M_S + M_p}$. On en déduit si $M_S \gg M_p$:

$$p = \frac{C^2}{G M_S}$$

Dans la suite, on prendra $\theta_0 = 0$.

a) Demi-grand axe a

Le demi-grand axe a est donné par : $2a = P_1P_2$ avec $SP_1 = \frac{p}{1+e}$ (P_1 correspond à $\theta = 0$) et $SP_2 = \frac{p}{1-e}$ (P_2 correspond à $\theta = \pi$, voir figure 31.12) donc

$$2a = \frac{p}{1+e} + \frac{p}{1-e} = \frac{2p}{1-e^2}$$

soit

$$a = \frac{p}{1-e^2}$$

b) Distance focale c

Il s'agit de la distance du centre de l'ellipse au centre de force ou foyer de l'ellipse.

soit

$$c = a - SP_1 = \frac{pe}{1-e^2}$$

c) Excentricité

Des deux relations donnant la distance focale et le demi-grand axe, on peut déduire l'excentricité de l'ellipse :

$$e = \frac{c}{a}$$

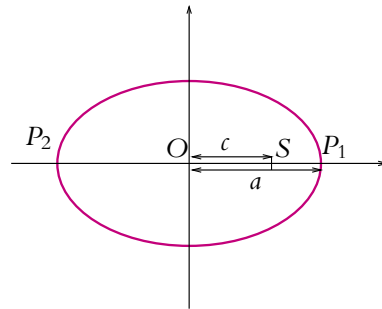


Figure 31.12 Trajectoire elliptique.

d) Demi petit axe b

Les relations générales sur les coniques donnent la relation :

$$a^2 = b^2 + c^2$$

donc

$$b^2 = a^2 - c^2 = \frac{p^2}{(1-e^2)^2} - \frac{p^2 e^2}{(1-e^2)^2} = \frac{p^2}{1-e^2}$$

On en déduit :

$$b = \frac{p}{\sqrt{1-e^2}} = a\sqrt{1-e^2} = \sqrt{ap}$$

5.4 Expression de l'énergie mécanique (MPSI, PCSI)

On a vu au paragraphe 4.3 que l'énergie mécanique et l'excentricité de la trajectoire sont reliées par l'équation :

$$Em = \frac{K^2}{2\mu C^2} (e^2 - 1)$$

De plus, on a établi que :

$$p = \mu \frac{C^2}{|K|}$$

On en déduit :

$$Em = \frac{|K|}{2p} (e^2 - 1)$$

Or on sait que $p = a(1 - e^2)$ donc :

$$Em = -\frac{|K|}{2a} \quad (31.2)$$

L'énergie mécanique ne dépend donc que du demi grand axe de l'ellipse. Ce résultat est à connaître.

L'expression obtenue est bien négative, ce qui est normal pour une trajectoire elliptique (qui correspond à un état lié).

On notera qu'on peut utiliser la relation obtenue dans le cas de la trajectoire circulaire et remplacer R par a le demi-grand axe de l'ellipse. Cette approche fournit un moyen pratique de retrouver l'expression (31.2) tout en n'en constituant pas une démonstration dans le cas général.

5.5 Période de révolution (MPSI, PCSI)

On a établi que pour un mouvement à force centrale, la vitesse aréolaire est une constante du mouvement. La période peut donc être obtenue comme le rapport entre l'aire totale de l'ellipse et de la vitesse aréolaire :

$$T = \frac{\text{Aire de l'ellipse}}{\text{Vitesse aréolaire}}$$

Or l'aire d'une ellipse est donnée par πab et la vitesse aréolaire par $\frac{C}{2}$, donc :

$$T = \frac{2\pi ab}{C}$$

Or

$$a = \frac{p}{1 - e^2} \quad \text{et} \quad b = \sqrt{ap}$$

D'autre part, $p = \frac{\mu C^2}{|K|}$ donc :

$$T = \frac{2\pi a \sqrt{ap}}{C} = \frac{2\pi a^{\frac{3}{2}}}{C} \sqrt{\frac{\mu C^2}{|K|}} = \frac{2\pi a^{\frac{3}{2}} \sqrt{\mu}}{|K|^{\frac{1}{2}}}$$

avec $|K| = GM_S M_P$ et $\mu = \frac{M_S M_P}{M_S + M_P}$. On en déduit :

$$T = \frac{2\pi a^{\frac{3}{2}}}{\sqrt{G(M_S + M_P)}}$$

et si $M_P \ll M_S$:

$$T = \frac{2\pi a^{\frac{3}{2}}}{\sqrt{GM_S}}$$

5.6 Troisième loi de Kepler (MPSI, PCSI)

De cette dernière relation, on déduit :

$$T^2 = \frac{4\pi^2 a^3}{GM_S}$$

soit :

$$\frac{T^2}{a^3} = \frac{4\pi^2}{GM_S} = \text{cte}$$

et on a établi la troisième loi de Kepler, la constante ne dépend pas de la planète considérée.

Si on ne peut pas écrire que $M_S \gg M_P$, on obtient :

$$\frac{T^2}{a^3} = \frac{4\pi^2}{G(M_S + M_P)}$$

qui n'est plus indépendant de M_P et donc de la planète.

Là encore, l'étude de la trajectoire circulaire permet d'obtenir rapidement la troisième loi de Kepler. On en déduit la relation pour la trajectoire elliptique en remplaçant R par a , ce qui constitue un moyen pratique de la retrouver et qu'on n'hésitera pas à employer.

5.7 Variation de la vitesse sur la trajectoire (MPSI, PCSI)

On a obtenu l'expression de l'énergie mécanique en fonction du demi grand axe :

$$Em = -\frac{GM_S M_P}{2a}$$

Or

$$Em = \frac{1}{2}\mu v^2 - \frac{GM_S M_P}{r}$$

avec

$$\mu = \frac{M_S M_P}{M_S + M_P} \simeq M_P$$

On en déduit l'expression de la vitesse le long de la trajectoire :

$$v^2 = \frac{2GM_S M_P}{2M_P} \left(\frac{2}{r} - \frac{1}{a} \right) = GM_S \left(\frac{2}{r} - \frac{1}{a} \right)$$

Or r varie entre $r_P = \frac{p}{1+e}$ et $r_A = \frac{p}{1-e}$. La vitesse étant une fonction décroissante de r , elle varie entre :

$$v(r_P) = (1+e) \sqrt{\frac{GM_S}{p}}$$

et

$$v(r_A) = (1-e) \sqrt{\frac{GM_S}{p}}$$

On remarque que la vitesse est maximale pour $r = r_P$, c'est-à-dire pour le point le plus proche ou *périgée*, et minimale pour $r = r_A$, à savoir pour le point le plus éloigné ou *apogée*.

6. Quelques remarques sur le mouvement des satellites (MPSI, PCSI)

6.1 Position du problème

Le mouvement des satellites terrestres peut être rapproché de ce qui vient d'être fait à propos des planètes. En effet, si on considère que le satellite est essentiellement soumis à l'attraction terrestre, on a à nouveau un problème à deux corps. On peut également supposer sans difficulté que la masse du satellite est très faible devant celle de la Terre. On assimilera par souci de simplicité le satellite à un point matériel et on travaillera dans le référentiel géocentrique défini au chapitre 27.

Tout ce qui a été obtenu pour le mouvement des planètes est donc transposable aux satellites. La question qui se pose est celle de la satellisation d'objets depuis la Terre. C'est à ce type de questions qu'on va s'intéresser maintenant. On fait l'hypothèse dans la suite que la trajectoire est elliptique. On n'a en effet aucun intérêt *a priori* à ne pas pouvoir récupérer le satellite et à ce qu'il parte à l'infini.

6.2 Vitesse circulaire ou première vitesse cosmique

Lorsqu'on détermine l'énergie mécanique du satellite de masse m , on peut le faire à l'instant $t = 0$ puisqu'il y a conservation de l'énergie mécanique au cours du mouvement :

$$Em = \frac{1}{2}mv_0^2 - \frac{GM_T m}{r_0}$$

en notant M_T la masse de la Terre, v_0 sa vitesse initiale et r_0 sa distance initiale au centre de la Terre.

On constate donc que l'énergie mécanique est une fonction croissante de la vitesse initiale : plus la vitesse initiale est importante, plus l'énergie mécanique est grande.

D'autre part, on a établi que l'énergie mécanique peut se mettre sous la forme :

$$Em = -\frac{GM_T m}{2a}$$

où a désigne le demi grand axe de l'ellipse. On en déduit donc que l'énergie mécanique est une fonction croissante de a .

On en déduit que la valeur du demi grand axe sera d'autant plus importante que la vitesse initiale sera grande.

Initialement, on choisit l'orbite basse la plus simple pour le satellite : une orbite circulaire de faible altitude donc de rayon proche de R_T . La vitesse du satellite est donnée par (Cf. paragraphe 5.2) :

$$v_0 = \sqrt{\frac{GM_T}{R_T}}$$

Cette vitesse est appelée première vitesse cosmique et est égale à la vitesse que doit initialement avoir un satellite à la surface de la Terre. Cette vitesse ne dépend pas de la masse du satellite.

Application numérique :

- $G = 6,67 \cdot 10^{-11} \text{ N.m}^2.\text{kg}^{-2}$, $M_T = 6 \cdot 10^{24} \text{ kg}$ et $R_T = 6400 \text{ km}$, on en déduit :

$$v_0 = \sqrt{\frac{6,67 \cdot 10^{-11} \cdot 6 \cdot 10^{24}}{6,4 \cdot 10^6}} = 7,9 \cdot 10^3 \text{ m.s}^{-1} = 28,5 \cdot 10^3 \text{ km.h}^{-1}$$

- en passant par le champ de pesanteur au sol :

$$mg_0 = \frac{mM_T G}{R_T^2}$$

on a :

$$g_0 = \frac{M_T G}{R_T^2} \simeq 10 \text{ m.s}^{-2}$$

et

$$v_0 = \sqrt{g_0 R_T}$$

Le raisonnement qui a été fait ici peut être obtenu à n'importe quelle altitude mais la première vitesse cosmique correspond par définition au cas où l'altitude h du satellite est faible devant le rayon terrestre R_T .

6.3 Vitesse de libération ou seconde vitesse cosmique

On a vu que dans le cas d'un potentiel newtonien attractif, ce qui est le cas envisagé ici, on peut avoir aussi bien un état libre pour lequel le point matériel part à l'infini qu'un état lié.

Dans le cas d'un satellite qui part à l'infini on dit qu'il se libère de l'attraction terrestre à laquelle il était soumis. Il s'agit d'un abus de langage : le satellite est toujours soumis à l'attraction terrestre même sur une trajectoire de type hyperbolique. La libération du satellite ne sera possible que si sa trajectoire n'est plus une ellipse mais une parabole ou une hyperbole. L'objet de ce paragraphe est de déterminer la vitesse nécessaire pour atteindre un état libre donc une trajectoire parabolique ou hyperbolique.

D'après la discussion qualitative sur la nature du mouvement, il faut que l'énergie mécanique soit nulle pour que l'état puisse devenir libre. Or on a :

$$E_m = \frac{1}{2}mv^2 - \frac{GmM_T}{R}$$

Si $E_m = 0$ alors :

$$v = \sqrt{\frac{2GM_T}{R}}$$

Si on suppose que le satellite est initialement à proximité de la surface de la Terre où on a $R \simeq R_T$, il faudra transmettre la vitesse dite *vitesse de libération*

$$v_l = \sqrt{\frac{2GM_T}{R_T}} = \sqrt{2}v_0$$

pour que le satellite puisse atteindre un état libre. Alors si $v > v_l$, alors $E_m > 0$ et le satellite garde une vitesse non nulle « à l'infini ».

Pour lancer un satellite sans qu'il atteigne sa vitesse de libération (ce qui reviendrait à le perdre !), il faut choisir une vitesse entre v_0 et $v_0\sqrt{2}$. La marge est donc étroite.

Application numérique :

$$v_l = 11,3 \text{ km.s}^{-1} = 40,7 \cdot 10^3 \text{ km.h}^{-1}$$

6.4 Troisième vitesse cosmique

Si la vitesse initiale v_i du satellite est supérieure à v_l alors, à l'infini¹, il lui reste une vitesse v_f calculable par l'application de la conservation de l'énergie mécanique :

$$E_m = \frac{1}{2}mv_i^2 - \frac{GM_Tm}{R_T} = \frac{1}{2}mv_i^2 - \frac{1}{2}mv_l^2 = \frac{1}{2}mv_f^2$$

¹ Cela correspond à sa « sortie » de l'attraction terrestre.

soit

$$v_f = \sqrt{v_i^2 - v_l^2}$$

Une fois libéré de l'attraction terrestre, on peut considérer que le satellite se trouve soumis dans le référentiel de Copernic à l'attraction solaire. Il est à une distance du Soleil assimilable à la distance Terre-Soleil et sa vitesse s'obtient par composition des vitesses en ajoutant la vitesse d'entraînement de la Terre. On a donc une nouvelle vitesse initiale :

$$v'_i = v_f + v_T$$

en notant v_T la vitesse de la Terre.

On appelle troisième vitesse cosmique v_{ls} la vitesse minimale dans le référentiel géocentrique nécessaire pour que le satellite se libère du système solaire. Par un raisonnement analogue à celui mené pour la seconde vitesse cosmique, on obtient la vitesse nécessaire dans le référentiel de Copernic :

$$v/\mathcal{R}_C = \sqrt{2}v_T$$

en assimilant la trajectoire de la Terre à un cercle. Donc v_{ls} est solution de :

$$v/\mathcal{R}_C = \sqrt{2}v_T = v_f + v_T = \sqrt{v_{ls}^2 - v_l^2} + v_T$$

$$\left(\sqrt{2} - 1\right)^2 v_T^2 = v_{ls}^2 - v_l^2$$

$$v_{ls} = \sqrt{v_l^2 + v_T^2 \left(\sqrt{2} - 1\right)^2}$$

Application numérique :

La distance Terre-Soleil vaut $d = 150.10^6$ km et la vitesse angulaire de révolution de la Terre par rapport au Soleil est de :

$$\omega = \frac{2\pi}{365,25 \times 24 \times 3600} = 2.10^{-7} \text{ rad.s}^{-1}$$

et

$$v_T = \omega d = 150.10^9 \times 2.10^{-7} = 29,9 \text{ km.s}^{-1} = 107,6.10^3 \text{ km.h}^{-1}$$

La troisième vitesse cosmique vaut donc :

$$v_{ls} = \sqrt{(11,3.10^3)^2 + (29,9.10^3)^2 \left(\sqrt{2} - 1\right)^2} = 16,8 \text{ km.s}^{-1} = 60,5.10^3 \text{ km.h}^{-1}$$

6.5 Rayon de la trajectoire d'un satellite géostationnaire

Un satellite est dit géostationnaire s'il a la même période de rotation T que la Terre sur elle-même, à savoir 24 h ou 86 400 s.

Pour que le satellite tourne à la même vitesse que la Terre, il est nécessaire que sa trajectoire soit circulaire et dans le plan équatorial. On a vu que sinon la vitesse variait en fonction de la distance au centre de la Terre. La vitesse du satellite est donnée (Cf. paragraphe 5.2) par :

$$v = \sqrt{\frac{GM_T}{R}}$$

en notant R le rayon de la trajectoire du satellite.

Or la vitesse angulaire du satellite vaut :

$$\omega = \frac{v}{R} = \frac{2\pi}{T}$$

soit

$$T = 2\pi \frac{R}{v} = 2\pi \sqrt{\frac{R^3}{GM_T}}$$

On en déduit le rayon d'une orbite géostationnaire de la Terre :

$$R = \sqrt[3]{\frac{GM_T T^2}{4\pi^2}}$$

Application numérique :

$$R = \sqrt[3]{\frac{6,67 \cdot 10^{-11} * 6 \cdot 10^{24} * 86\,400^2}{4\pi^2}} = 42\,300 \text{ km} = 6,6R_T$$

L'altitude correspond à $h = R - R_T = 5,6R_T = 36\,000 \text{ km}$ environ.

Les données numériques suivantes sont utilisées dans de nombreux exercices :

- Masse de la Terre : $M_T = 5,97 \cdot 10^{24}$ kg.
- Rayon de la Terre : $R_T = 6,38 \cdot 10^6$ m.
- Masse du Soleil : $M_S = 1,99 \cdot 10^{30}$ kg.
- Constante universelle de gravitation : $G = 6,67 \cdot 10^{-11}$ N.m².kg⁻².
- Valeur du champ de pesanteur au niveau du sol : $g = 9,81$ m.s⁻².
- Unité astronomique (distance moyenne Terre-Soleil) : 1 u.a. = $1,5 \cdot 10^{11}$ m.

A. Applications directes du cours

1. Explosion d'une comète

Une comète de période $T = 770$ ans s'est approchée du soleil à $d = 7,75 \cdot 10^{-3}$ u.a..

1. On suppose la trajectoire elliptique. Déterminer le demi-grand axe a , l'excentricité e , le paramètre p de l'ellipse, et les vitesses à l'aphélie (v_A) et au périhélie (v_P). A-t-on eu raison de faire l'hypothèse d'une trajectoire elliptique ?

2. À son périhélie, elle a explosé en deux morceaux de masse m_1 et m_2 , partis respectivement avec les vitesses $\vec{v}_1$ et $\vec{v}_2$ dans deux directions faisant des angles respectifs $\alpha_1 = 20^\circ$ et $\alpha_2 = 50^\circ$ avec la direction de la vitesse initiale $\vec{v}_P$ (figure 31.13). La quantité de mouvement totale se conserve durant l'explosion et l'on observe que $\|\vec{v}_1\| \simeq \|\vec{v}_2\|$.

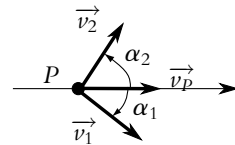


Figure 31.13

a) En utilisant la conservation de la quantité de mouvement, déterminer le rapport des masses m_1/m_2 puis la norme des vitesses des fragments.

b) Calculer l'énergie mécanique par unité de masse pour chaque fragment. En déduire le type de la nouvelle trajectoire.

2. Trajectoire d'un satellite et vitesse initiale

Un satellite de la Terre est sur une orbite circulaire de rayon r_0 . En un point A on lui communique subitement une vitesse $\vec{v}_A$ perpendiculaire au rayon vecteur $\vec{r}_0$ pour le faire passer sur une nouvelle trajectoire.

1. Exprimer la vitesse v_c sur l'orbite circulaire en fonction de M_T (masse de la Terre), r_0 et G .

2. a) Que peut-on dire du point A pour la nouvelle trajectoire ? En déduire la relation entre p (paramètre), r_0 et e (excentricité) pour la nouvelle trajectoire (deux relations sont possibles suivant le type de trajectoire).

b) Exprimer l'excentricité de la nouvelle trajectoire en fonction de v_A et v_c .

c) Calculer e pour $v_A = v_c$, $v_A = v_c/3$ et $v_A = 3v_c$. Donner le type de trajectoire pour chaque cas.

3. Retombée d'un satellite sur Terre

Un satellite est sur une orbite circulaire de rayon r_0 . En raison de dysfonctionnements on souhaite le faire retomber sur Terre dans une zone non habitée. Pour cela, à partir de l'orbite initiale circulaire on réduit sa vitesse brutalement au point A pour qu'il retombe sur Terre après avoir tourné d'un angle de 90° en suivant une trajectoire elliptique.

1. Le point A sur l'orbite circulaire étant l'apogée de la nouvelle trajectoire elliptique, déterminer la relation entre le paramètre de l'ellipse p et R_T puis la relation entre l'excentricité e , r_0 et R_T .
2. Établir la relation entre l'énergie mécanique et le demi grand axe de l'ellipse a . En déduire l'expression de l'énergie mécanique en fonction de G , m (masse du satellite), M_T , R_T et r_0 .
3. Quelle est la vitesse à donner au satellite en A pour qu'il s'écrase sur Terre à l'endroit souhaité ?

4. Étude du mouvement d'une planète par la méthode du vecteur excentricité, d'après ENSTIM 2007

On veut étudier le mouvement d'une planète P , assimilée à un point matériel dans le champ de gravitation d'une étoile de masse M_e , de centre O , considérée comme ponctuelle et fixe. La planète de masse M_p est située à une distance $r = OP$ de O . On considère un référentiel lié à l'étoile comme galiléen.

1. a) Exprimer la force exercée par l'étoile sur la planète en fonction de masses M_p et M_e , r , G la constante de gravitation universelle et $\vec{u}_r = \vec{OP}/r$.
- b) Justifier précisément que le mouvement est plan. Préciser ce plan. On notera $(\vec{u}_r, \vec{u}_\theta)$ la base de projection dans ce plan et $\vec{u}_z$ un vecteur unitaire suivant la direction du moment cinétique en O , $\vec{L} = L\vec{u}_z$. Déterminer l'expression de L en fonction de M_p , r et θ .

On rappelle que l'équation polaire d'une ellipse est $r(\theta) = \frac{p}{1 + e \cos \theta}$. On définit le vecteur excentricité :

$$\vec{e} = -\frac{L}{GM_e M_p} \vec{v} + \vec{u}_\theta$$

où $\vec{v}$ est la vitesse de la planète.

2. a) Montrer que $\vec{e}$ est un vecteur constant.

En fait $\vec{e}$ est un vecteur orthogonal au grand axe de l'ellipse (figure 31.14).

- b) En calculant le produit scalaire $\vec{e} \cdot \vec{u}_\theta$, montrer que l'on retrouve l'équation de l'ellipse.

- c) Que vaut le module de $\vec{e}$? Préciser p en fonction de G , M_e , M_p et L .

- d) Préciser la valeur de l'excentricité dans le cas d'un mouvement circulaire.

- e) Dans le cas du mouvement circulaire, préciser la valeur de L en fonction de R (rayon du cercle), v_c (vitesse circulaire) et M_p . Déterminer l'expression de la vitesse v_c en fonction de R , G et M_e à l'aide du vecteur excentricité.

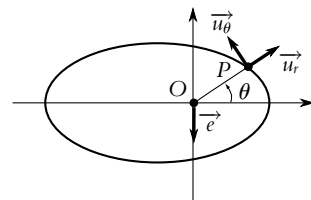


Figure 31.14

5. Chute d'un point sur une planète

Un point matériel P de masse m tombe sur une planète de masse M et de rayon R . On suppose que la planète possède une répartition de masse à symétrie sphérique et que le point P est abandonné sans vitesse initiale d'une hauteur $h = 4R$. On fait également l'hypothèse que $M \gg m$.

Calculer le temps de chute.

6. Satellite terrestre, d'après ICNA 2007

Dans les séries de réponses proposées, il peut y avoir plusieurs réponses justes.

Un satellite (S) est assimilé à une masse ponctuelle $m = 2$ tonnes. La Terre (T) est supposée être sphérique de rayon R_T , de masse M_T et de répartition massique uniforme. On note :

- G la constante universelle de gravitation ;
- $\mathcal{R}(T, x, y, z)$ le référentiel galiléen géocentrique ;
- $\vec{g}$ le champ de pesanteur terrestre subit par S dans $\mathcal{R}$;
- $B(\vec{u}_r, \vec{u}_\theta, \vec{u}_\varphi)$ la base sphérique orthonormée directe liée à $\mathcal{R}$ et associée à S .

1. Dans $\mathcal{R}$, la somme des forces extérieures $\sum \vec{F}_{\text{ext}}(S)/\mathcal{R}$ est :

- A) Radiale.
- B) Orthoradiale.
- C) Centrale.
- D) Centrifuge.

2. Soit $\vec{L}_T(S)/\mathcal{R}$ le moment cinétique de S en T par rapport à $\mathcal{R}$. Le point S est repéré par ses coordonnées sphériques dans $\mathcal{R}$.

- A) $\vec{L}_T(S)/\mathcal{R}$ est dirigé suivant $\vec{u}_\theta$.
- B) $\vec{L}_T(S)/\mathcal{R}$ est indépendant de la masse du satellite.
- C) $\vec{L}_T(S)/\mathcal{R}$ est nul.
- D) $\vec{L}_T(S)/\mathcal{R}$ est une constante vectorielle.

3. Le mouvement de S autour de la Terre est maintenant supposé être plan, uniforme et circulaire à une altitude h autour de la Terre. Le point S est repéré par ses coordonnées polaires (ρ, φ) dans $\mathcal{R}$, et on définit la base $B_{p0}(\vec{u}_\rho, \vec{u}_\varphi)$ de $\mathcal{R}$, permettant de définir la position de S dans le plan du mouvement.

Soit $\vec{a}/\mathcal{R}(S)$ et $\vec{v}/\mathcal{R}(S)$, respectivement l'accélération et la vitesse de S dans $\mathcal{R}$.

$$\text{A) } \vec{a}/\mathcal{R}(S) = \left[\frac{d\vec{v}/\mathcal{R}(S)}{dt} \right]_{\mathcal{R}} \quad \text{B) } \|\vec{a}/\mathcal{R}(S)\| = \left[\frac{d\|\vec{v}/\mathcal{R}(S)\|}{dt} \right]_{\mathcal{R}}$$

$$\text{C) } \vec{a}/\mathcal{R}(S) = \frac{\|\vec{v}/\mathcal{R}(S)\|^2}{R_T + h} \vec{u}_\rho \quad \vec{a}/\mathcal{R}(S) = -\frac{\|\vec{v}/\mathcal{R}(S)\|^2}{R_T + h} \vec{u}_\rho$$

4. La norme de la vitesse $\vec{v}_{/\mathcal{R}}(S)$ vérifie :

$$\begin{array}{ll} \text{A) } \|\vec{v}_{/\mathcal{R}}(S)\| = \sqrt{\frac{GM_T^2}{R_T + h}} & \text{B) } \|\vec{v}_{/\mathcal{R}}(S)\| = \sqrt{\frac{Gm^2}{R_T}} \\ \text{C) } \|\vec{v}_{/\mathcal{R}}(S)\| = \sqrt{\frac{GM_T m}{R_T + h}} & \text{D) } \|\vec{v}_{/\mathcal{R}}(S)\| = \sqrt{\frac{GM_T m}{R_T}} \end{array}$$

5. La période T_0 de révolution du satellite est :

- A) Indépendante de m .
- B) Indépendante de h .
- C) Proportionnelle à $h^{3/2}$.
- D) Proportionnelle à $R_T^{3/2}$.

6. Sachant que $G = 6,67 \cdot 10^{-11}$ SI, $M_T = 6 \cdot 10^{27}$ g, $R_T = 6400$ km et que la période de révolution est à peu près 24 h, l'altitude pour laquelle le satellite est géostationnaire est environ de :

- A) 36 000 m.
- B) 36 000 km.
- C) 417 000 m.
- D) 417 000 km.

7. On note E_c , E_p et E respectivement les énergies cinétique, potentielle et mécanique du satellite géostationnaire.

$$\text{A) } E_c = \frac{1}{2} \frac{Gm^2 M_T}{R_T + h} \quad \text{B) } E_c = 9400 \text{ MJ}$$

$$\text{C) } E_c = \frac{1}{2} \frac{Gm M_T^2}{R_T + h} \quad \text{D) } E_c = 1,9 \cdot 10^{13} \text{ J}$$

8.

$$\text{A) } E_p = -\frac{Gm^2 M_T}{R_T + h} \quad \text{B) } E_p = -18,9 \text{ GJ}$$

$$\text{C) } E_p = -\frac{Gm M_T^2}{R_T + h} \quad \text{D) } E_p = -3,8 \cdot 10^{13} \text{ J}$$

9.

$$\text{A) } E = -\frac{1}{2} \frac{Gm^2 M_T}{R_T + h} \quad \text{B) } E = -9,5 \text{ GJ}$$

$$\text{C) } E = -E_c \quad \text{D) } E = -1,9 \cdot 10^{13} \text{ J}$$

10. Un satellite S' de masse m' a une trajectoire elliptique (demi-grand axe a), la Terre étant un des foyers de l'ellipse.

Au moment où S' passe par son périégée :

- A) La distance entre S' et T est minimale.
- B) La distance entre S' et T est maximale.
- C) La vitesse de S' est minimale.
- D) La vitesse de S' est maximale.

11. La période du satellite S' étant de 7 jours :

- A) $a = 23.10^3$ km B) $a = 155.10^3$ km
- C) $a = 119.10^3$ km D) $a = 785.10^3$ km

B. Exercices et problèmes

1. Déviation d'un météorite

Un météorite a, très loin de la Terre, une vitesse $v_0 = v_0 \vec{u}_x$, portée par une droite Δ située à la distance b du centre de la Terre. On note m la masse du météorite, M celle de la Terre, R le rayon de la Terre et G la constante de la gravitation universelle.

1. Déterminer la distance minimale Terre-Météorite (notée d).
2. Déterminer la valeur minimale $b_{\min}$ de b pour que le météorite ne rencontre pas la Terre.
3. Dans le cas où $b > b_{\min}$, déterminer l'angle de déviation φ du météorite.
4. On donne $v_0 = 11 \text{ km.s}^{-1}$. Calculer $b_{\min}$. Calculer φ pour $b = 1,5b_{\min}$.

2. Trajectoire quasi-circulaire d'un satellite - Freinage par l'atmosphère

On étudie le mouvement d'un satellite artificiel de la Terre dans le référentiel géocentrique supposé galiléen. On néglige les autres interactions que la force de gravitation entre la Terre et le satellite. On note M_T la masse de la Terre, R_T son rayon, m la masse du satellite supposée petite devant M_T et G la constante de gravitation universelle. On note T_0 la période de révolution du satellite.

1. Établir la conservation du moment cinétique du satellite par rapport à la Terre.
2. En déduire que le mouvement du satellite est plan.
3. Montrer que cela permet de définir une constante des aires C dont on donnera l'expression.
4. On suppose que le satellite est en orbite circulaire autour de la Terre. Montrer que son mouvement est uniforme.
5. Établir l'expression de la vitesse v du satellite en fonction de G , M_T et r le rayon de l'orbite du satellite.
6. Déterminer l'expression de l'énergie cinétique E_c du satellite en fonction de G , M_T , m et r .
7. Même question pour l'énergie potentielle E_p du satellite. Donner la relation entre E_c et E_p .
8. En déduire l'expression de l'énergie mécanique E_m et les relations de E_m avec E_p et E_c .

9. Ce satellite subit une force de frottement fluide proportionnelle à sa masse et au carré de sa vitesse. On note α le coefficient de frottement correspondant. Cette force reste suffisamment faible pour que la trajectoire soit quasi-circulaire. Dans ces conditions, les relations entre les différentes énergies restent valables et la variation Δr du rayon de la trajectoire est faible devant le rayon r . Déterminer ΔE_p la variation d'énergie potentielle sur une révolution du satellite consécutive à une variation Δr du rayon de son orbite.
10. En déduire que l'énergie mécanique du satellite diminue sur une révolution d'une quantité ΔE_m qu'on explicitera.
11. Calculer sur une révolution le travail W_f de la force de frottement en fonction de α , m , v et T_0 .
12. En déduire que $\Delta r = -4\pi\alpha r^2$.
13. Quel est par conséquent l'effet de la force de frottement sur le rayon de la trajectoire et la vitesse du satellite ?
14. Retrouver la troisième loi de Kepler dans le cas d'une trajectoire circulaire.
15. En supposant que la variation du rayon sur une révolution peut s'identifier à la dérivée du rayon par rapport au temps, montrer que $\sqrt{r} = \sqrt{r_0} + Kt$ où K est une constante à exprimer en fonction de G , M_T et α .

3. Comète à trajectoire parabolique

La comète Arend-Roland est une comète à trajectoire d'excentricité e estimée à 1,0002. On assimilera la trajectoire à une parabole d'excentricité $e = 1$. La comète est passée à son périhélie, le 8 avril 1957, à $r_p = 0,316$ u.a. du soleil.

1. a) Calculer le paramètre p de la trajectoire.
- b) Exprimer le paramètre p de la conique en fonction de la constante des aires C , de G et la masse du soleil M_S . En déduire la vitesse de la comète au périhélie.
2. Montrer que l'on peut exprimer la vitesse angulaire $\dot{\theta}$ en fonction de la constante des aires C et du paramètre p de la parabole par :

$$\frac{d\theta}{dt} = \dot{\theta} = \frac{C(1 + \cos \theta)^2}{p^2}$$

En quel point est prise l'origine $\theta = 0$?

3. a) On note t_0 l'instant de passage au périhélie P . Montrer alors qu'on obtient le temps de passage pour le point d'angle α avec :

$$t(\alpha) = t - t_0 + \int_0^\alpha \frac{p^2}{C} \frac{d\theta}{(1 + \cos \theta)^2}$$

- b) Pour calculer cette intégrale, on effectue le changement de variable $x = \tan \frac{\theta}{2}$. En déduire l'expression :

$$t = t_0 + \frac{p^2}{2C} \left(X + \frac{X^3}{3} \right)$$

où $X = \tan \frac{\alpha}{2}$.

c) La comète a été découverte par les astronomes belges Sylvain Arend et Georges Roland le 8 novembre 1956. À quelle distance du soleil se trouvait-elle ?

4. Mise en orbite d'une sonde spatiale

On souhaite mettre en orbite une sonde spatiale (s) autour de Mars. Ce véhicule possède au point A la vitesse $\vec{v}_A$ et présente un « paramètre d'impact » b (figure 31.15).

Données :

- $v_A = 2,5 \text{ km.s}^{-1}$;
- rayon de Mars : $r_M = 3397 \text{ km}$;
- masse de Mars : $M_M = 6,39.10^{23} \text{ kg}$.

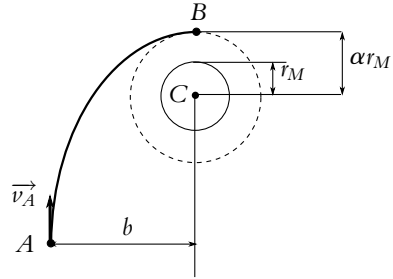


Figure 31.15

1. Au point A , on suppose que le véhicule est assez éloigné de Mars pour pouvoir négliger l'énergie gravitationnelle. En déduire l'expression de l'énergie mécanique et la nature de la trajectoire de (S).

2. Montrer que le moment cinétique se conserve et définir la constante des aires C . Calculer C en fonction de v_A et b .

Sachant que la trajectoire d'approche est tangente au cercle de rayon αr_M en B , calculer v_B (vitesse en C) en fonction de v_A , r_M , α et b .

3. Exprimer le paramètre d'impact b , en fonction de r_M , v_A , M_M et α . Application numérique pour $\alpha = 3$.

4. Déterminer la distance minimale b_m pour que le véhicule évite la surface de Mars.

5. Déterminer la vitesse v_c d'un objet sur l'orbite circulaire de rayon $3r_M$ ainsi que sa période de révolution en fonction des données. Effectuer l'application numérique.

6. Au point C (avec $\alpha = 3$), on veut que le véhicule passe sur l'orbite circulaire de rayon $3r_M$. Déterminer la variation de vitesse Δv_c à communiquer au véhicule en fonction des données. Application numérique.

5. Mise en orbite d'un satellite, orbite de transfert

La mise en orbite d'un satellite géostationnaire à l'aide d'un lanceur type Ariane est réalisée de la façon suivante :

1. Le lanceur injecte le satellite au périégée d'une orbite de transfert elliptique d'inclinaison i correspondant à la latitude de l'aire de lancement.
2. Le satellite se positionne sur la trajectoire géostationnaire à l'aide de son moteur d'apogée (lorsqu'il est à l'apogée de l'orbite de transfert).

La période de rotation de la Terre est $T = 86\,164 \text{ s}$ et la masse du satellite $m = 700 \text{ kg}$.

On donne la formule de Binet pour l'accélération :

$$\vec{a} = -C^2 u^2 \left(\frac{d^2 u}{d\theta^2} + u \right) \vec{u}_r$$

où $u = 1/r$ et C est la constante des aires.

1. L'orbite géostationnaire est celle d'un satellite restant toujours à la verticale d'un même point du globe.
 - a) Calculer la vitesse angulaire Ω_g d'un satellite sur l'orbite géostationnaire. Quel est le référentiel d'étude ?
 - b) Déterminer l'expression du rayon de cette orbite R_g ainsi que son altitude h_g par rapport à la Terre, puis de la vitesse v_g du satellite sur cette orbite. Montrer que cette orbite est forcément dans le plan de l'équateur. Effectuer les applications numériques pour R_g , h_g et v_g .
2. Le lanceur injecte le satellite au périégée P de l'orbite de transfert (qui est une ellipse) à une altitude $h_0 = 200$ km et avec une vitesse telle que l'apogée A de l'orbite de transfert soit sur l'orbite géostationnaire.
 - a) En déduire le demi grand-axe a de l'ellipse de transfert ainsi que son excentricité e et son paramètre p .
 - b) En utilisant la formule de Binet, établir la relation donnant la constante des aires C en fonction de M_T , G et p . En déduire la vitesse v_0 à donner au satellite lors de son injection au périégée de l'orbite de transfert, ainsi que la vitesse v_1 à l'apogée. Faire les applications numériques.
 - c) Calculer le temps que passe le satellite sur l'orbite de transfert.

6. Satellite artificiel, d'après Agrégation

1. Lancement d'un satellite

On étudie le lancement d'un satellite artificiel à partir d'un point O de la surface terrestre.

- a) Établir l'expression de la vitesse du point O dans le référentiel géocentrique $\mathcal{R}_J$ (assimilé ici à un référentiel galiléen) en fonction de la vitesse de rotation de la Terre autour de l'axe de ses pôles Ω , du rayon terrestre R_T et de la latitude du lieu λ .
- b) En déduire les conditions les plus favorables pour le lancement du satellite. Parmi les trois champs de tirs suivants, lequel choisir de préférence ?
 - Baïkonour au Kazakhstan : $\lambda = 46^\circ$;
 - Cap Canaveral aux USA : $\lambda = 28,5^\circ$;
 - Kourou en Guyane française : $\lambda = 5,23^\circ$.
- c) Établir l'expression de l'énergie potentielle de gravitation du système Terre-satellite en fonction de l'altitude z du satellite par rapport au sol. On prend pour référence une énergie potentielle nulle à l'infini. En déduire l'expression de l'énergie mécanique du satellite sur sa base de lancement dans le référentiel géocentrique.
- d) On appelle ici vitesse de libération v_l , la vitesse verticale minimale qu'il faut communiquer initialement au satellite par rapport au sol, pour qu'il puisse se libérer de l'attraction terrestre. Donner l'expression de v_l . Calculer sa valeur numérique dans le cas où le satellite est lancé de la base de Kourou.

2. Satellite artificiel en orbite

On considère un satellite artificiel de masse m en mouvement circulaire autour de la Terre.

- a) Montrer que le mouvement du satellite est uniforme. Établir l'expression de la vitesse du satellite en fonction de son altitude ainsi que la troisième loi de Kepler liant la période de rotation T du satellite au rayon r de sa trajectoire.
- b) Calculer le rayon de l'orbite d'un satellite géostationnaire et définir son plan de révolution.
- c) Quelle énergie cinétique minimale faut-il communiquer au satellite pour qu'il échappe à l'attraction terrestre s'il est initialement en orbite circulaire autour de la Terre à l'altitude z ? A.N. : $z = 36\,000\text{ km}$; $m = 6\text{ tonnes}$.
- d) Soit un satellite d'énergie initiale E_{m0} . Son orbite est relativement basse et il subit donc les frottements des couches hautes de l'atmosphère. Il s'ensuit que l'énergie mécanique du satellite varie selon la loi : $E_m = E_{m0}(1 + bt)$, b étant un coefficient constant positif. On suppose que la trajectoire reste approximativement circulaire.
- Préciser le signe de E_{m0} . Établir l'expression du rayon r et de la vitesse v du satellite en fonction du temps. Comparer les évolutions de r et de v ainsi que celles des énergies potentielle et cinétique. Que devient l'énergie perdue ?

7. Satellite Hipparcos, d'après ENSTIM 2000

Le satellite Hipparcos lancé le 8 Août 1987 était constitué principalement d'un télescope de 30 cm de diamètre. Celui-ci a permis d'établir un catalogue des positions, distances et éclats de plus de 118 000 étoiles avec une précision jamais atteinte. Ce satellite devait être placé sur une orbite géostationnaire à une altitude $H = 36\,000\text{ km}$. Un problème de mise à feu du moteur d'apogée a laissé Hipparcos sur son orbite de transfert son altitude variant entre h et H . Après utilisation des moteurs de positionnement, l'altitude minimale a été portée à $h = 500\text{ km}$. Une programmation du satellite a permis de s'affranchir des problèmes liés à cette orbite. Au cours d'une révolution, il passe dans la ceinture de Van Allen. On supposera que cette ceinture est comprise entre deux sphères de rayon $r_1 = 8\,400\text{ km}$, $r_2 = 28\,000\text{ km}$ et de centre celui de la terre. La ceinture de Van Allen est constituée de particules chargées piégées dans le champ magnétique terrestre. Ces particules aveuglent les détecteurs d'Hipparcos interrompant les mesures des positions des étoiles. Il est cependant utilisable à 65%. On assimile la terre à une sphère de centre O , de rayon $R = 6\,400\text{ km}$ et de masse M et le satellite à un point matériel (S, m) . On suppose le référentiel géocentrique $\mathcal{R}$ galiléen. La période de rotation de la terre dans ce référentiel appelée jour sidéral vaut $T = 86\,164\text{ s}$.

1. a) Quelle est la nature de la trajectoire d'Hipparcos ?
 - b) Déterminer les expressions de l'excentricité e et du paramètre de l'ellipse p en fonction de h , H et R . Application numérique.
 - c) Exprimer et calculer le demi-grand axe a de la trajectoire.
 - d) Démontrer la troisième loi de Képler.
 - e) Dédire de la question précédente, la relation entre la période de rotation de la Terre T et l'altitude de l'orbite géostationnaire H .
 - f) Exprimer la période T_h de révolution d'Hipparcos en fonction de T , R , H et h . Calculer T_h en heure.
2. a) Déterminer les valeurs numériques des angles θ_1 et θ_2 d'entrée et de sortie de la ceinture de Van Allen du satellite. On donnera les valeurs comprises entre 0° et 180° .

- b) Représenter sur un schéma clair la trajectoire du satellite et l'aire S_b balayée par $\overrightarrow{OS}$ lors d'un passage dans la ceinture de Van Allen. Pour la question suivante, on prendra une valeur approchée de $S_b = 200 \cdot 10^6 \text{ km}^2$.
- c) Déterminer le rapport $\rho = t_0/T_h$ en fonction de S_b et S_e (surface de l'ellipse) où t_0 est la durée totale d'inactivité d'Hipparcos sur une période. Application numérique.

8. Changement d'orbite d'un satellite

On souhaite transférer un satellite depuis une orbite circulaire rasante de rayon R_T autour de la Terre sur son orbite géostationnaire de rayon R_G . On fera l'étude dans le référentiel géocentrique dans lequel la Terre tourne sur elle-même à la vitesse angulaire Ω . On néglige les autres interactions que la force de gravitation entre la Terre et le satellite. On note O le centre de la Terre, $M_T = 5,98 \cdot 10^{24} \text{ kg}$ sa masse, $R_T = 6,37 \cdot 10^6 \text{ m}$ son rayon, $m = 1,5 \text{ t}$ la masse du satellite supposée petite devant M_T et G la constante de gravitation universelle.

- Établir que la trajectoire du satellite géostationnaire est forcément dans le plan équatorial.
- En déduire que les trois orbites appartiennent à un même plan à préciser.
- Déterminer la vitesse v_B du satellite sur son orbite basse avant son transfert sur son orbite géostationnaire.
- Donner sa valeur numérique.
- Exprimer le rayon de la trajectoire géostationnaire.
- Donner sa valeur numérique.
- Préciser l'altitude de l'orbite géostationnaire.
- Calculer la valeur numérique de la vitesse du satellite sur son orbite géostationnaire.
- Le transfert du satellite de son orbite basse à son orbite géostationnaire s'effectue de la manière suivante : on communique au satellite une brusque variation de vitesse en un point P de sa trajectoire basse en éjectant des gaz pendant un intervalle de temps très court dans le sens opposé à la vitesse du satellite. Il suit alors une orbite elliptique et lorsque sa trajectoire croise la droite OP au point A , on lui communique un supplément de vitesse pour le stabiliser sur l'orbite géostationnaire. Établir une relation entre les vitesses aux points A et P et les distances $r_A = OA$ et $r_P = OP$.
- En utilisant la conservation de l'énergie mécanique sur la trajectoire elliptique, établir l'expression de l'énergie mécanique en fonction de G , M_T , m et a le demi grand axe de l'ellipse.
- Donner la valeur numérique de l'énergie mécanique sur l'ellipse.
- Établir l'expression de la vitesse du satellite sur la trajectoire elliptique en fonction de R_G , R_T , r , G et M_T .
- Donner la valeur de la variation de vitesse qu'il faut imposer en P .
- En déduire la variation de l'énergie mécanique en P .
- Donner la valeur de la variation de vitesse qu'il faut imposer en A .
- En déduire la variation de l'énergie mécanique en A .

17. Déterminer la durée de ce transfert.

9. Satellites sélénostationnaires, d'après CCP DEUG 2006

La Lune L de masse M décrit une orbite circulaire de rayon $a = 384,4 \cdot 10^3$ km et centrée sur la Terre T de masse kM avec $k = 81$. Le référentiel terrestre $\mathcal{R}_0$ est supposé galiléen. On note $\mathcal{R}_1$ le référentiel lié à la Lune et obtenu par une rotation d'angle θ_1 autour de l'axe Tz_0 , on suppose que les axes Tz_0 et Lz_1 sont parallèles. On néglige l'influence de tous les autres astres que la Terre et la Lune.

1. En appliquant le principe fondamental de la dynamique à la Lune, exprimer la vitesse de rotation ω_1 de la Lune autour de la Terre en fonction de k , G la constante de gravitation universelle, M et a .

2. Où doit-on placer une particule de masse m pour qu'elle se trouve au point d'équigravité du système Terre - Lune ? On exprimera la distance x de ce point à la Lune en fonction de a et k .

3. Donner la valeur numérique de x .

4. La Lune tourne sur elle-même avec la même période de révolution T_0 qu'autour de la Terre. Que peut-on en déduire ?

5. On cherche à obtenir un satellite de masse m sélénostationnaire c'est-à-dire immobile par rapport à la surface de la Lune. On note $\mathcal{R}_2$ le référentiel lié au satellite et obtenu par une rotation d'angle θ_2 autour de Lz_1 .

Établir la relation simple entre ω_1 et ω_2 les vitesses de rotation de la Lune autour de la Terre et du satellite autour de la Lune.

6. En considérant que $\mathcal{R}$ est galiléen et qu'on puisse négliger l'attraction terrestre sur le satellite, déterminer le rayon de son orbite en fonction de a et k .

7. Donner sa valeur numérique.

8. Que peut-on conclure de ce résultat ?

9. Dans la suite, on tient compte de l'attraction terrestre. Il existe sur l'axe Terre - Lune deux positions tels qu'un satellite placé à cet endroit soit sélénostationnaire. Déterminer les équations vérifiées par x_1 et x_2 en notant $x_1 a$ et $x_2 a$ les distances de ces points à la Lune.

10. En supposant que x_1 et x_2 sont petits devant 1, résoudre ces équations.

11. Calculer les valeurs numériques des distances de ces points à la Lune.

12. Brusquement la Lune perd sa vitesse orbitale et se met à chuter vers la Terre en suivant une trajectoire elliptique. Déterminer la durée de la chute en supposant qu'elle est inférieure à une période de révolution de la Lune autour de la Terre.

13. Donner sa valeur numérique en supposant que $T_0 = 27$ jours 7 h 43 min 11s.

Quelques aspects de la mécanique terrestre (PCSI)

32

L'objet de ce chapitre est d'une part de préciser les notions de champs de gravitation et de pesanteur, d'autre part d'étudier quelques problèmes liés à ces notions comme les phénomènes de marée et la déviation vers l'est d'un objet en chute libre.

1. Interaction gravitationnelle

1.1 Force gravitationnelle

Dans le chapitre sur les principes de la dynamique newtonienne, on a introduit les divers types d'interaction dont l'interaction gravitationnelle, qui est traitée dans ce chapitre.

On rappelle que deux points matériels M_1 et M_2 de masses respectives m_1 et m_2 exercent l'un sur l'autre une force attractive, proportionnelle aux masses m_1 et m_2 et inversement proportionnelle au carré de la distance entre M_1 et M_2 . Cette force dite de Newton s'écrit mathématiquement :

$$\overrightarrow{f_{M_1 \rightarrow M_2}} = -G \frac{m_1 m_2}{(M_1 M_2)^3} \overrightarrow{M_1 M_2}$$

en notant $G = 6,672 \cdot 10^{-11} \text{ N.m}^2.\text{kg}^{-2}$ la constante de gravitation.

On peut formuler quelques remarques sur cette force :

1. Elle vérifie le principe des actions réciproques : M_1 exerce une force sur M_2 et réciproquement M_2 exerce une force sur M_1 . Le caractère ponctuel des masses permet d'en déduire que la force s'exerce le long de $M_1 M_2$. La forme de son expression mathématique garantit que ces deux forces sont opposées.
2. Tout corps qui en attire d'autres est en même temps attiré par ces derniers. Il est donc à la fois actif et passif.

3. Cette force définit la notion de masse gravitationnelle comme un scalaire caractéristique de la quantité de matière considérée. Le choix de la constante G impose l'unité de masse comme étant le kilogramme dans les unités du système international. On a établi expérimentalement l'identité de cette masse gravitationnelle à la masse inertielle intervenant dans le principe fondamental de la dynamique. Cela permet donc de confondre les deux notions et de ne parler que d'une seule masse qu'elle soit gravitationnelle ou inertielle.

1.2 Champ gravitationnel

On peut introduire la notion de champ gravitationnel en considérant l'action subie par une masse fixe du fait de la présence en un point variable d'une autre masse. Ceci est facilement réalisable à partir de l'expression de la force gravitationnelle :

$$\overrightarrow{f_{M_1 \rightarrow M_2}} = m_2 \left(-G \frac{m_1}{(M_1 M_2)^3} \overrightarrow{M_1 M_2} \right) = m_2 \overrightarrow{A(M_2)}$$

On appelle champ gravitationnel créé en M_2 par la masse m_1 placée en M_1 la quantité :

$$\overrightarrow{A(M_2)} = -G \frac{m_1}{(M_1 M_2)^3} \overrightarrow{M_1 M_2}$$

Ce champ fait partie d'une classe plus générale de champs : les champs newtoniens ou champs en $\frac{1}{r^2}$ comme le champ électrostatique qui sera étudié dans le cours d'électromagnétisme. Comme on le verra en électromagnétisme, la notion de champ gravitationnel qui vient d'être définie pour une masse M_2 peut se généraliser à n'importe quel type de distribution de masses dans l'espace. On ne développera pas plus cette notion de champ gravitationnel, on reviendra sur ce point dans le cours d'électromagnétisme où on obtiendra différentes propriétés par analogie avec le champ électrostatique. La seule chose qu'on retiendra ici est qu'une masse quelconque m placée en M est soumise à une force $m \overrightarrow{A(M)}$ du fait de la présence d'autres masses, et que $\overrightarrow{A(M)}$ est le champ gravitationnel en M .

2. Champ de pesanteur terrestre

Dans la suite, on se limite à l'étude du champ de gravitation terrestre, c'est-à-dire au champ gravitationnel créé par la Terre.

2.1 Référentiel terrestre

Le référentiel approprié à cette étude est le référentiel terrestre défini comme le référentiel lié à la Terre. Son origine est le centre de la Terre et ses axes suivent la Terre dans son mouvement.

On a vu lors de l'étude de la recherche d'un référentiel galiléen que ce référentiel n'est pas rigoureusement galiléen du fait d'une part de la rotation de la Terre sur elle-même, et d'autre part de son mouvement de translation elliptique autour du Soleil. On devra donc tenir compte des éventuelles forces d'inertie pour interpréter certains phénomènes.

On rappelle que le référentiel terrestre est en rotation par rapport au référentiel géocentrique et que le référentiel géocentrique est défini par son origine au centre de la Terre, ses axes étant les mêmes que celui du référentiel de Copernic (ils ne suivent donc pas la Terre dans sa rotation sur elle-même).

Le référentiel géocentrique est lui-même en translation par rapport au référentiel de Copernic dont l'origine est le centre de masse du système solaire et les axes dirigés vers trois étoiles fixes très éloignées. On peut également rappeler un référentiel proche de celui de Copernic : le référentiel de Kepler dont l'origine est le centre du Soleil au lieu d'être le centre de masse du système solaire. Dans la suite, on considère le référentiel de Copernic comme galiléen.

Le référentiel terrestre qu'on va utiliser est donc en translation et en rotation par rapport au référentiel de Copernic. Il n'est pas galiléen.

Pour le mouvement d'un point aux environs de la surface de la Terre, on utilise un référentiel terrestre local. On le définit à partir des notations du schéma suivant :

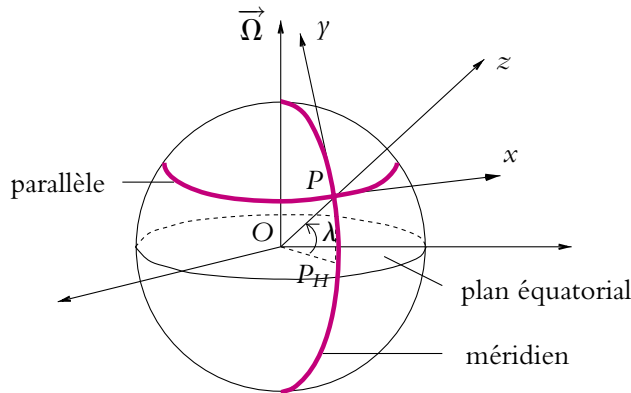


Figure 32.1 Référentiel terrestre : base locale.

Soit un point P à la surface de la Terre. On note P_H sa projection dans le plan équatorial ou plan passant par le centre O de la Terre et perpendiculaire à l'axe de rotation de la Terre sur elle-même.

On définit la *latitude* λ du point P comme l'angle entre OP_H et OP . On appelle *parallèle* le cercle à la surface de la Terre dont tous les points ont même latitude.

Un demi-cercle de centre O et passant par les pôles est appelé *méridien* ; sa position est définie par la *longitude*.

On choisit l'axe Pz selon la verticale locale au point P , l'axe Px le long du parallèle passant par P , dirigé vers l'Est, et l'axe Py le long du méridien passant par P , dirigé vers le Nord.

Dans la suite du paragraphe, on considère le référentiel géocentrique comme galiléen.

2.2 Bilan des forces dans le référentiel terrestre

On considère une particule ponctuelle de masse m .

Elle est soumise :

1. à la force de gravitation de la Terre $\vec{A}_T$,
2. à la force d'inertie d'entraînement $\vec{f}_{ie} = -m\vec{a}_e$, où $\vec{a}_e$ est l'accélération du point coïncidant dans le référentiel géocentrique,
3. à la force d'inertie de Coriolis $\vec{f}_{ic} = -m\vec{a}_c = -2m\vec{\Omega} \wedge \vec{v}$, en notant $\vec{\Omega}$ le vecteur rotation de la Terre sur elle-même et $\vec{v}$ la vitesse du point par rapport au référentiel terrestre.

On néglige, dans la suite, toute autre force et on ne s'intéresse qu'aux aspects gravitationnels et au caractère non galiléen du référentiel terrestre.

2.3 Définition du poids d'un corps

Pour définir le poids d'un corps, on imagine celui-ci accroché à un fil. Il n'est soumis à aucune autre interaction que l'attraction gravitationnelle terrestre et la tension du fil.

Le poids de ce corps de masse m , noté $\vec{P}$, est la force opposée à celle qui le maintient en équilibre (c'est-à-dire à la tension du fil) dans le référentiel terrestre en considérant le référentiel géocentrique comme galiléen.

On va chercher à exprimer cette force compte tenu de ce qui a été dit dans les paragraphes précédents.

- La masse m est en équilibre, donc sa vitesse par rapport au référentiel terrestre est nulle : $\vec{v} = \vec{0}$, ce qui implique la nullité de la force de Coriolis.
- Les forces s'exerçant sur la masse m sont donc la force d'inertie d'entraînement $\vec{f}_{ie} = -m\vec{a}_e$, la force de gravitation $m\vec{A}_T$ et la force $\vec{F}$ qui la maintient en équilibre (tension du fil).

On en déduit que la force $\vec{F}$ qui maintient la masse en équilibre vérifie :

$$\vec{F} + m\vec{A}_T - m\vec{a}_e = \vec{0}$$

Donc, par définition du poids, on déduit :

$$\vec{P} = -\vec{F} = m\vec{A}_T - m\vec{a}_e$$

Le poids est donc la somme de la force d'inertie d'entraînement du référentiel terrestre par rapport au référentiel géocentrique et de la force de gravitation de la Terre.

On a souvent l'habitude d'écrire sous la forme :

$$\vec{P} = m \vec{g}$$

en appelant *champ de pesanteur* et en notant $\vec{g}$ la quantité :

$$\vec{g} = \vec{A}_T - \vec{a}_c$$

2.4 Expression de la force d'inertie d'entraînement

L'expression générale de l'accélération d'entraînement¹ est :

$$\vec{a}_c(M) = \vec{a}(O) + \vec{\Omega} \wedge (\vec{\Omega} \wedge \vec{OM}) + \frac{d\vec{\Omega}}{dt} \wedge \vec{OM} = \vec{a}(O) - \Omega^2 \vec{HM} + \frac{d\vec{\Omega}}{dt} \wedge \vec{OM}$$

en notant H la projection de M sur l'axe de rotation, c'est-à-dire sur l'axe des pôles.

On néglige les variations de la rotation de la Terre sur elle-même. Cela revient à ne pas considérer les phénomènes de précession des équinoxes, les phénomènes dus aux oscillations angulaires de l'axe de rotation² de direction et les phénomènes liés au freinage de la rotation de la Terre sur elle-même. Ne pas considérer les variations

de la rotation de la Terre sur elle-même revient à négliger le terme $\frac{d\vec{\Omega}}{dt} \wedge \vec{OM}$ dans l'expression de l'accélération d'entraînement du référentiel terrestre par rapport au référentiel géocentrique.

Le point O est fixe dans le référentiel géocentrique : le terme $\vec{a}(O)$ est nul.

Finalement on n'a que le terme de force centrifuge dans la force d'inertie d'entraînement :

$$\vec{f}_{ie} = m \Omega^2 \vec{HM}$$

On notera que le référentiel géocentrique n'est pas tout à fait galiléen. On verra ultérieurement quelles sont les conséquences de cet aspect.

► **Remarque** Si on fait directement l'hypothèse que la vitesse de rotation de la Terre sur elle-même est constante, le point coïncidant est animé d'un mouvement circulaire de

¹ C'est la seule fois où l'expression générale peut être utile.

² Cela signifie que l'axe des pôles se déplace légèrement et très lentement.

centre H donc son accélération est :

$$\vec{a}_e(M) = -\Omega^2 \overrightarrow{HM}$$

On obtient l'expression de l'accélération d'entraînement sans utiliser la formule générale rappelée ci-dessus.

2.5 Expression du champ de pesanteur terrestre

Le champ de pesanteur terrestre a pour expression :

$$\begin{aligned} \vec{g} &= \vec{A}_T(M) - \vec{a}_e(M) \\ &= \vec{A}_T(M) + \Omega^2 \overrightarrow{HM} \end{aligned}$$

On remarque que le champ de pesanteur n'est pas dirigé selon la droite reliant M au centre de la Terre. Il faut déterminer l'importance relative des différents termes pour calculer l'écart entre la verticale (donnée par la direction de $\vec{g}$) et la droite reliant M au centre de la Terre.

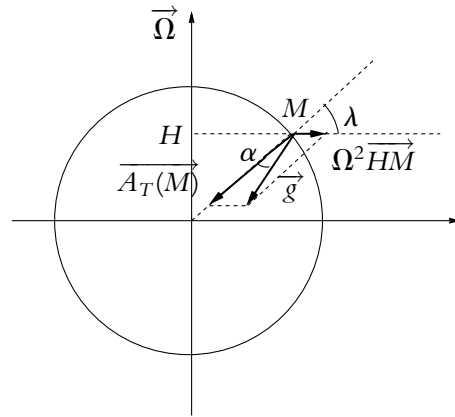


Figure 32.2 Orientation du champ de pesanteur.

2.6 Importance relative des différents termes

a) Terme d'attraction terrestre

Il s'agit du terme $\vec{A}_T(M)$ dirigé du point M vers le centre O de la Terre et dont le module vaut :

$$A_T(M) = \frac{GM_T}{R_T^2}$$

en notant M_T et R_T respectivement la masse et le rayon de la Terre et en supposant que la Terre est à symétrie sphérique.

Application numérique :

$$M_T = 6.10^{24} \text{ kg}, R_T = 6,4.10^6 \text{ m et } G = 6,67.10^{-11} \text{ m}^3.\text{kg}^{-1}.\text{s}^{-2}.$$

D'où $A_T(M) = 9,77 \text{ m.s}^{-2}$.

Il s'agit du terme prépondérant comme le montreront les applications numériques ultérieures.

b) Terme de la force centrifuge

Il s'agit du terme $\Omega^2 \overrightarrow{HM}$ qui est proportionnel au carré de la vitesse de rotation³ de la Terre sur elle-même :

$$\Omega = \frac{2\pi}{24 \times 3600} = 7,27 \cdot 10^{-5} \text{ rad.s}^{-1}$$

et proportionnel à la distance du point à l'axe de rotation.

On peut exprimer cette distance en fonction du rayon de la Terre et de la latitude λ du point considéré. La latitude est définie comme l'angle que fait la droite passant par le point considéré et le centre de la Terre, et le plan équatorial (Cf. figure ci-contre). On a donc :

$$HM = R_T \cos \lambda$$

Ce terme dépend fortement de la latitude à laquelle on se place :

- aux pôles $HM = 0$ et le terme centrifuge est nul,
- à l'équateur $HM = R_T$ et le terme centrifuge vaut $0,034 \text{ m.s}^{-2}$,
- à la latitude de 45° ⁴,

$$HM = R_T \cos 45^\circ = 4,5 \cdot 10^6 \text{ m}$$

et la contribution centrifuge vaut $0,024 \text{ m.s}^{-2}$.

Il varie donc entre 0 et $0,034 \text{ m.s}^{-2}$. L'ordre de grandeur de ce terme montre qu'il s'agit d'une correction par rapport au terme précédent lié à l'attraction terrestre.

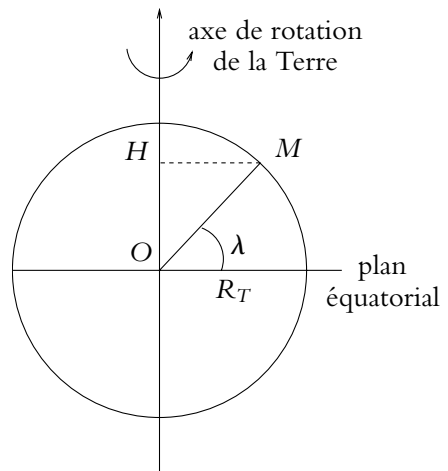


Figure 32.3 Expression de la force centrifuge.

c) Inclinaison du champ de pesanteur

L'angle α entre la direction du champ de pesanteur et la droite OM est donné en un point M par la relation :

$$\tan \alpha = \frac{a_e(M) \sin \lambda}{A_T(M) + a_e(M) \cos \lambda} \simeq \frac{a_e(M) \sin \lambda}{A_T(M)}$$

³ On rappelle que la Terre effectue un tour sur elle-même en un jour.

⁴ Il s'agit approximativement de la latitude de la France, le parallèle 45° Nord passe en France à hauteur de Valence.

compte tenu des ordres de grandeurs des termes du dénominateur.

Pour un point à la latitude de 45° :

$$\tan \alpha = 1,74 \cdot 10^{-3} \quad \text{donc} \quad \alpha = 1,74 \cdot 10^{-3} \text{ rad} = 5'58''.$$

Cet angle est très faible, ce qui permet de considérer en première approximation que le champ de pesanteur est dirigé vers le centre de la Terre.

2.7 Principe fondamental de la dynamique dans le référentiel terrestre

Si on considère que le référentiel géocentrique est galiléen, les forces qui s'exercent sur un point matériel m dans le référentiel terrestre sont :

- le poids $m\vec{g}$ avec l'expression du champ de pesanteur $\vec{g}$ qui vient d'être définie,
- la force d'inertie de Coriolis $-2m\vec{\Omega} \wedge \vec{v}_{/\mathcal{R}_T}(M)$ en notant $\vec{v}_{/\mathcal{R}_T}(M)$ la vitesse du point M par rapport au référentiel terrestre,
- d'éventuelles autres forces $\vec{f}_a$.

Le principe fondamental de la dynamique s'écrit donc :

$$m\vec{a}_{/\mathcal{R}_T}(M) = m\vec{g} - 2m\vec{\Omega} \wedge \vec{v}_{/\mathcal{R}_T}(M) + \vec{f}_a$$

3. Terme de marée

On s'intéresse maintenant à l'influence gravitationnelle d'autres astres sur le mouvement d'un point matériel dans le champ de pesanteur. On doit alors tenir compte des conséquences du fait que le référentiel géocentrique n'est pas galiléen. On considère ici que le référentiel galiléen de référence est le référentiel de Copernic.

3.1 Définition

On tient également compte dans la force de gravitation de l'influence des autres astres. On se limite dans la suite à un seul astre autre que la Terre. Cette hypothèse n'est en aucun cas une perte de généralités : il suffira de sommer les contributions dues aux éventuels autres astres du fait de la linéarité des équations.

La force de gravitation se décompose donc en un terme dû à l'attraction terrestre, $\vec{A}_T(M)$, et un terme lié à l'autre astre $\vec{A}_a(M)$:

$$\vec{A}(M) = \vec{A}_T(M) + \vec{A}_a(M)$$

L'expression de l'accélération d'entraînement devient maintenant :

$$\vec{a}_e(M) = \vec{a}_{/\mathcal{R}_C}(O) - \Omega^2 \overrightarrow{HM}$$

en notant $\vec{a}_{/\mathcal{R}_C}(O)$ l'accélération du centre de la Terre dans le référentiel de Copernic. Il n'y a que ce terme à ajouter par rapport au cas précédent où on considérait le référentiel géocentrique comme galiléen puisque ce dernier est en translation par rapport au référentiel de Copernic. On néglige toujours les variations de la vitesse de rotation de la Terre sur elle-même.

On en déduit l'expression du principe fondamental de la dynamique dans le référentiel terrestre :

$$m\vec{a}_{/\mathcal{R}_T}(M) = m\vec{A}_T(M) + m\vec{A}_a(M) - m\vec{a}_{/\mathcal{R}_C}(O) + m\Omega^2\overline{HM} - 2m\vec{\Omega} \wedge \vec{v}_{/\mathcal{R}_T}(M)$$

soit en introduisant le champ de pesanteur $\vec{g}$ (qui regroupe $\vec{A}_T(M)$ et le terme centrifuge) :

$$m\vec{a}_{/\mathcal{R}_T}(M) = m\vec{g} + m\vec{A}_a(M) - m\vec{a}_{/\mathcal{R}_C}(O) - 2m\vec{\Omega} \wedge \vec{v}_{/\mathcal{R}_T}(M)$$

Or, on peut établir que $\vec{a}_{/\mathcal{R}_C}(O) = \vec{A}_a(O)$: il suffit d'appliquer le principe fondamental de la dynamique à la Terre dans le référentiel de Copernic supposé galiléen. La Terre n'est soumise dans ce cas qu'à l'interaction gravitationnelle des autres astres. D'autre part, il faut supposer que les astres peuvent être assimilés à un point, à savoir qu'ils présentent une symétrie sphérique et qu'ils sont très éloignés les uns des autres de manière à pouvoir négliger leur rayon devant la distance entre leur centre et celui des autres astres. On a alors

$$M_T\vec{a}_{/\mathcal{R}_C}(O) = M_T\vec{A}_a(O)$$

d'où le résultat. On justifiera le fait que le champ de gravitation créé par un astre à symétrie sphérique puisse être assimilé à celui d'un point matériel au centre de l'astre et portant la totalité de la masse dans le cours d'électromagnétisme (dans le chapitre consacré aux calculs de champs).

Finalement, le principe fondamental de la dynamique appliqué à un point matériel de masse m (en l'absence d'autres phénomènes d'interaction que celui de gravitation) s'exprime par :

$$m\vec{a}_{/\mathcal{R}_T}(M) = m\vec{g} - 2m\vec{\Omega} \wedge \vec{v}_{/\mathcal{R}_T}(M) + m\left(\vec{A}_a(M) - \vec{A}_a(O)\right)$$

Le terme : $\vec{A}_a(M) - \vec{A}_a(O)$ est le terme de marée. Il va permettre d'analyser le phénomène du même nom.

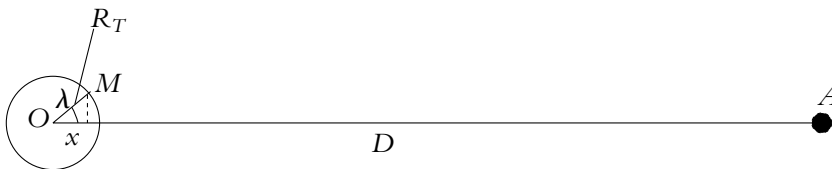


Figure 32.4 Expression du terme de marée.

On notera M_A la masse du second astre et D la distance de son centre à celui de la Terre.

On suppose en outre que tous les astres sont à symétrie sphérique.

Dans ces conditions, le terme de marée est la différence de :

$$\vec{A}_a(M) = -GM_A \frac{\vec{AM}}{AM^3} \quad \text{et} \quad \vec{A}_a(O) = -GM_A \frac{\vec{AO}}{AO^3}$$

3.2 Ordres de grandeur

On suppose que $R_T \ll D$. On en déduit que :

$$\left\| \frac{\vec{AM}}{AM^3} - \frac{\vec{AO}}{AO^3} \right\| \simeq \frac{\|\vec{AM} - \vec{AO}\|}{D^3} \simeq \frac{R_T}{D^3}$$

L'ordre de grandeur du terme de marée est donc celui de :

$$\frac{GM_A R_T}{D^3}$$

Les valeurs numériques pour le Soleil, la Lune, Vénus et Jupiter sont réunies dans le tableau suivant :

Astres	Masse (kg)	Distance à la Terre (km)	Terme de marée (m.s ⁻²)
Soleil	1,99.10 ³⁰	1,50.10 ¹¹	2,5.10 ⁻⁷
Lune	7,35.10 ²²	3,84.10 ⁸	5,5.10 ⁻⁷
Vénus	4,92.10 ²⁴	4,2.10 ¹⁰	2,8.10 ⁻¹¹
Jupiter	1,91.10 ²⁷	6,3.10 ¹¹	3,3.10 ⁻¹²

Plusieurs enseignements peuvent être tirés de ces valeurs numériques :

- le terme de marée peut être considéré comme une correction du terme d'attraction terrestre et du terme lié à la rotation de la Terre sur elle-même.
- les effets de la Lune et du Soleil sont comparables, alors que ceux liés aux autres planètes du système solaire comme Vénus ou Jupiter peuvent être négligés devant ceux de la Lune et du Soleil.

Tout ce qui est exposé dans la suite n'est pas exigible dans le cadre du programme mais explique les effets visibles du terme de marée à la surface de la Terre.

3.3 Déformations de la surface de la Terre

a) Étude qualitative

Le terme de marée est la différence de deux vecteurs :

$$\overrightarrow{A_a(M)} = -GM_A \frac{\overrightarrow{AM}}{AM^3} \quad \text{et} \quad \overrightarrow{A_a(O)} = -GM_A \frac{\overrightarrow{AO}}{AO^3}$$

On le représente graphiquement en différents points de la surface de la Terre :

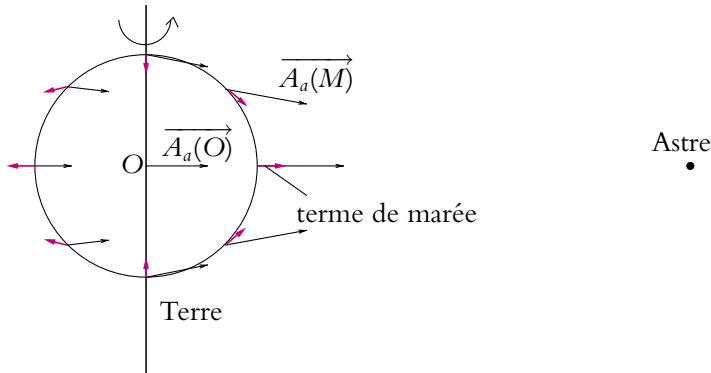


Figure 32.5 Termes de marée à la surface de la Terre.

Les effets sont symétriques par rapport à l'axe des pôles : le terme de marée est le même, dirigé vers l'extérieur, pour deux points situés sur l'équateur par exemple. C'est à l'équateur qu'il est maximal. En revanche, au niveau des pôles, il est minimal et dirigé vers le centre de la Terre.

On observe donc une déformation de la surface de la Terre en fonction de la latitude avec un renflement au niveau de l'équateur connu sous la dénomination de « bourrelet ». La déformation diminue ensuite régulièrement de l'équateur aux pôles (Cf. figure 32.6).

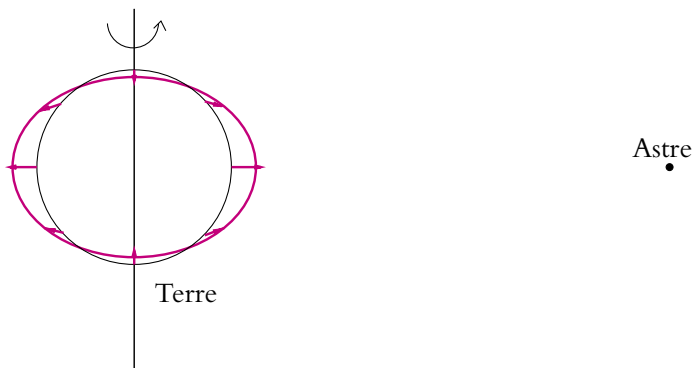


Figure 32.6 Déformations de la Terre par phénomène de marée.

b) Étude quantitative

On cherche à exprimer plus précisément le terme de marée.

Pour cela, on va utiliser une base de projection cartésienne dont l'axe Ox sera la droite OA et on adopte les notations suivantes :

$$\overrightarrow{OA} = D\vec{u}_x \quad \text{et} \quad \overrightarrow{MA} = (D - x)\vec{u}_x - y\vec{u}_y - z\vec{u}_z$$

avec $x \ll D$, $y \ll D$ et $z \ll D$ ou encore $OM = \sqrt{x^2 + y^2 + z^2} = R_T \ll D$.

Compte tenu de ce rapport entre le rayon de la Terre et la distance à l'astre, on va effectuer un développement limité de $\frac{1}{MA^3}$:

$$\begin{aligned} \frac{1}{MA^3} &= ((D - x)^2 + y^2 + z^2)^{-\frac{3}{2}} = \frac{1}{D^3} \left(1 - 2\frac{x}{D} + \frac{x^2 + y^2 + z^2}{D^2} \right)^{-\frac{3}{2}} \\ &\simeq \frac{1}{D^3} \left(1 + 3\frac{x}{D} \right) \end{aligned}$$

en se limitant au premier ordre en $\frac{x}{D}$.

On en déduit :

$$\vec{A}_a(M) - \vec{A}_a(O) = GM_A \left(\frac{\overrightarrow{MA}}{D^3} + 3\frac{x\overrightarrow{MA}}{D^4} - \frac{\overrightarrow{OA}}{D^3} \right) = \frac{GM_A}{D^3} \left(\overrightarrow{MO} + 3\frac{x}{D}\overrightarrow{MA} \right)$$

au premier ordre en $\frac{x}{D}$.

Le premier terme s'ajoute dans la même direction que la contribution due à l'attraction terrestre tandis que le second dépend, tout comme le terme lié à la rotation de la Terre sur elle-même, de la latitude du point considéré. Le signe de cette deuxième contribution est celui de x et est positif pour la demi-sphère la plus proche de l'astre et négatif pour l'autre demi-sphère.

En coordonnées sphériques, en utilisant la latitude λ au lieu du classique angle θ , on obtient :

$$x = R_T \cos \lambda \cos \varphi$$

La contribution dépend fortement de la latitude :

- aux pôles où $\lambda = \frac{\pi}{2}$ et $\cos \lambda = 0$, la contribution est minimale au premier ordre car $x = 0$ et le second terme est nul,
- à l'équateur où $\lambda = 0$ et $\cos \lambda = 1$, les effets de la marée sont maximaux.

3.4 Existence des marées semi-diurnes

Le schéma ci-dessous indique le sens du terme de marée pour une latitude donnée : le seul paramètre variable est l'angle α entre la direction Ox et celle de HM dans un plan parallèle au plan équatorial. On constate donc que le terme de marée a tendance à déformer la surface de la Terre.

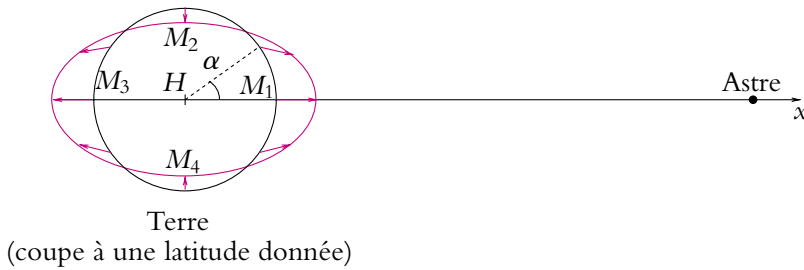


Figure 32.7 Déformations de la Terre à une latitude donnée par phénomène de marée.

Ces faits expliquent l'existence de deux marées quotidiennes. La Terre tournant en 24 h sur elle-même, un point de l'équateur passe une fois en position M_1 et une fois en position M_3 où la déformation est maximale. De même, il passe une fois en M_2 et une fois en M_4 où la déformation est minimale. On aura donc un phénomène d'oscillations entre une déformation minimale et une déformation maximale avec une période de 2 oscillations en 24 h.

La déformation maximale correspond à la marée dite haute et la déformation minimale à la marée dite basse.

La révolution de la Lune et le fait que la période de rotation de la Terre ne soit pas exactement de 24 heures entraînent un décalage de 50 minutes par jour pour les marées semi-diurnes.

3.5 Marées de mortes et de vives eaux

Jusqu'ici, on ne s'est intéressé à l'influence que d'un seul astre. Or les ordres de grandeurs ont montré qu'il était au moins nécessaire de considérer les effets conjoints du Soleil et de la Lune. Les échelles ne sont respectées pour aucune des figures de ce paragraphe.

Lorsque les trois astres Terre, Lune et Soleil sont alignés, les effets s'ajoutent. On obtient l'une des deux situations suivantes :

- à la pleine Lune :



Figure 32.8 Pleine Lune et marées.

- à la nouvelle Lune :



Figure 32.9 Nouvelle Lune et marées.

Dans les deux cas, on obtient une amplitude importante des marées : on parle de marées de « vives eaux ».

A contrario, lorsque les axes Terre-Lune et Terre-Soleil sont perpendiculaires, les amplitudes de marée s'annulent presque du fait que la Lune et le Soleil donnent des amplitudes comparables en ordre de grandeur et que lorsqu'un astre introduit une faible déformation, l'autre en introduit une plus grande et réciproquement. On obtient l'une des deux situations correspondant aux premier et dernier quartier de lune :



Figure 32.10 Marées de mortes eaux (au premier quartier).

On parle alors de marées de « mortes eaux » du fait des faibles variations d'amplitude.

3.6 Marées d'équinoxe

Au cours d'une année, le Soleil évolue dans le plan de l'écliptique qui fait un angle de $23^{\circ}27'$ avec le plan équatorial. De ce fait, l'influence du Soleil sur les marées évolue au cours d'une année.

Le Soleil traverse le plan équatorial aux équinoxes de printemps et d'automne (le 20 mars et le 22 septembre). Les effets des marées sont alors les plus importants : l'angle entre plan de l'écliptique et plan équatorial n'intervient pas. Ceci explique l'existence des marées d'équinoxe qui sont fortes.

À l'opposé, les 21 juin et 21 décembre, le Soleil est dans sa position la plus éloignée du plan équatorial et les marées sont les plus faibles.

3.7 Phénomène de résonance

L'amplitude des marées varie d'un point à un autre du globe en fonction de la profondeur des fonds marins. En effet, en certains endroits, la profondeur est telle que les oscillations propres des masses de fluide ont la même période que celle des marées.

Il y a alors un phénomène de résonance qui se produit : l'amplitude des marées est amplifiée. C'est par exemple ce qui se produit dans la baie du Mont-Saint-Michel. On ne s'intéresse ici qu'à un modèle statique des marées, ce qui exclut tous les effets liés à la forme des côtes, à la mer...

3.8 Marées, anneaux planétaires et limite de Roche

Ce qui vient d'être établi à propos des marées permet de fournir une explication à l'existence des anneaux planétaires de Saturne, Jupiter ou Uranus. En effet, les forces de marée ont tendance à disloquer l'astre sur lequel elles s'exercent. Une explication possible de la séparation de la masse d'un astre au niveau de l'équateur et de la formation d'anneaux planétaires serait l'existence de forces de marée plus importantes que les forces de cohésion de l'astre.

Pour les mêmes raisons, un objet passant à proximité d'un astre pour lequel la force de marée est importante peut voir la cohésion de sa structure brisée par la force de marée. Il existe une distance limite à l'astre appelée *limite de Roche* en dessous de laquelle ce phénomène se produit. Cela fournit une autre explication possible de l'existence des anneaux planétaires qui pourraient donc également être constitués de débris d'astres qui se seraient trop approchés de la planète.

On peut noter que l'effet de marée et son effet de dislocation a permis de réfuter la formation des planètes à partir de globules de matière issus du Soleil. Ces derniers seraient disloqués dès leur formation à la surface du Soleil et ne pourraient donc donner naissance à des planètes.

4. Exemple de l'influence de la force de Coriolis : déviation vers l'est

Dans ce paragraphe, on suppose que le référentiel géocentrique est galiléen et on montre à titre d'exercice un exemple de l'influence de la force de Coriolis. On s'intéresse à la chute libre d'un point matériel dans le champ de pesanteur en tenant compte du caractère non galiléen du référentiel terrestre et en particulier à l'influence de la force de Coriolis. On néglige les forces de frottement dans cet exemple.

Pour tenir compte du caractère non galiléen du référentiel terrestre, il faut considérer :

- la force d'inertie d'entraînement : celle-ci a été incluse par définition dans le champ de pesanteur et elle a déjà été prise en compte. On rappelle que cette force d'inertie d'entraînement explique partiellement les variations du champ de pesanteur à la surface de la Terre.
- la force d'inertie de Coriolis dont l'expression est :

$$-2m\vec{\Omega} \wedge \vec{v}$$

où $\vec{v}$ est la vitesse du point dans le référentiel terrestre.

On s'intéresse ici aux effets de cette force, puisque le poids prend en compte la partie de la force d'inertie due à la rotation du référentiel terrestre et on négligera le phénomène de marée.

4.1 Équation du mouvement

Le principe fondamental de la dynamique s'écrit en tenant compte de la force de Coriolis et en négligeant les forces de frottement :

$$m \vec{a} = m \vec{g} - 2m \vec{\Omega} \wedge \vec{v}$$

On se reportera à la figure définissant le référentiel terrestre du début du chapitre pour le choix de la base. On rappelle qu'on prend l'axe Oz comme la verticale du lieu, l'axe Ox le long du parallèle, vers l'Est, et l'axe Oy le long du méridien, vers le Nord.

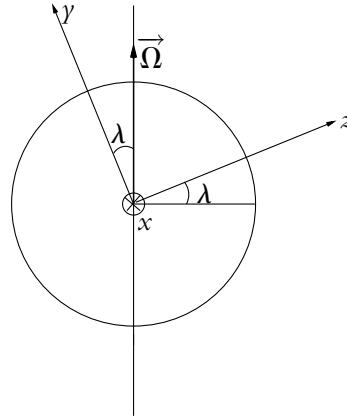


Figure 32.11 Coupe de la Terre dans un plan méridien.

4.2 Ordre de grandeur de la force de Coriolis

La Terre tourne sur elle-même en environ 24 h. Sa vitesse de rotation autour de son axe est donc :

$$\Omega = 7,27 \cdot 10^{-5} \text{ rad.s}^{-1}$$

Pour que cette force ait le même ordre de grandeur que le poids, il faut que l'ordre de grandeur de la vitesse soit :

$$v \simeq \frac{g}{2\Omega} \simeq 10^5 \text{ m.s}^{-1}$$

Cette vitesse est rarement atteinte par un point matériel usuel.

Pour un point matériel de vitesse 100 m.s^{-1} (ce qui est déjà une vitesse importante), l'importance relative des deux termes est :

$$\left| \frac{2\Omega v}{g} \right| \simeq 0,15 \%$$

Par conséquent, on peut en conclure que la force de Coriolis n'apporte qu'un terme correctif et on va la traiter comme une perturbation au premier ordre de la force de pesanteur.

4.3 Déviation vers l'Est

Étant donné que la force de Coriolis peut être considérée comme une perturbation par rapport au poids, on prend dans l'expression de la force de Coriolis l'expression de la vitesse non perturbée, c'est-à-dire sans tenir compte de la force de Coriolis, soit :

$$\vec{v} = \vec{v}_0 + \vec{g}t$$

Si la vitesse initiale est nulle (cas de la chute d'un corps lâché sans vitesse initiale), on a donc :

$$\vec{f}_{ic} = -2m\vec{\Omega} \wedge \vec{g}t$$

soit en projetant sur $Oxyz$ le principe fondamental de la dynamique :

$$\begin{cases} \ddot{x} = 2m\Omega \cos \lambda gt \\ \ddot{y} = 0 \\ \ddot{z} = -g \end{cases}$$

L'équation en x s'intègre en :

$$\dot{x} = \Omega \cos \lambda gt^2$$

puis en :

$$x = \frac{\Omega}{3} \cos \lambda gt^3$$

Compte tenu du choix de l'orientation des axes, on observe donc un déplacement dans le sens des x croissants à savoir une déviation vers l'Est. On peut noter que le sens de la déviation ne dépend pas de la latitude, il est donc le même dans les deux hémisphères, mais que son amplitude est fonction du cosinus de cette dernière.

Avec une vitesse initiale non nulle, un point matériel pourra être dévié vers l'Ouest car la force de Coriolis n'a pas la même direction pour une vitesse ascendante et descendante.

4.4 Ordre de grandeur de la déviation

On va faire l'application numérique avec les valeurs suivantes : $h = 50$ m, $\lambda = 45^\circ$ et on prendra pour $g = 9,8 \text{ m.s}^{-2}$.

Le temps de chute est solution de :

$$-\frac{1}{2}gT^2 + h = 0$$

soit

$$T = \sqrt{\frac{2h}{g}} = \sqrt{\frac{2 \times 50}{9,8}} = 3,19 \text{ s}$$

La déviation vaut alors :

$$\frac{\Omega}{3} \cos \lambda g T^3 = \frac{7,27 \cdot 10^{-5}}{3} \cos 45^\circ \times 9,8 \times 3,19^3 = 5,5 \text{ mm}$$

Cette valeur est faible, la force d'inertie de Coriolis est bien un terme correctif.

On peut citer l'expérience réalisée par Camille Flammarion (1842-1925) en 1903 à partir de la coupole du Panthéon à Paris d'une hauteur de 67 m et à une latitude de $48^\circ 51'$. On obtient théoriquement une durée de chute de 3,70 s et une déviation de 8,0 mm. Flammarion a déterminé expérimentalement une déviation de 7,6 mm, valeur proche de celle calculée ci-dessus.

4.5 Déviation Nord-Sud

À cet ordre d'approximation, la trajectoire reste dans le plan xOz . Pour mettre en évidence des effets d'ordre supérieur, il faut garder l'expression complète de la vitesse dans la force d'inertie de Coriolis, à savoir $\vec{v} = \dot{x}\vec{u}_x + \dot{y}\vec{u}_y + \dot{z}\vec{u}_z$. Le principe fondamental de la dynamique projeté sur Oy donne alors :

$$\ddot{y} = -2\Omega \sin \lambda \dot{x} = -2\Omega^2 \sin \lambda \cos \lambda t^2 = -\Omega^2 \sin (2\lambda) t^2$$

soit par intégration :

$$\dot{y} = -\Omega^2 \sin (2\lambda) \frac{t^3}{3}$$

et

$$y = -\Omega^2 \sin (2\lambda) \frac{t^4}{12}$$

Dans l'hémisphère Nord, λ est compris entre 0 et $\frac{\pi}{2}$ donc $\sin 2\lambda > 0$ et on aura une déviation vers le Sud.

Dans l'hémisphère Sud, λ est compris entre $-\frac{\pi}{2}$ et 0 donc $\sin 2\lambda < 0$ et on aura une déviation vers le Nord.

L'ordre de grandeur de cette déviation est encore plus faible : avec les mêmes valeurs que précédemment

$$y = -\frac{(7,27 \cdot 10^{-5})^2 \times \sin(2 \times 45^\circ) 3,19^4}{12} = -0,14 \text{ } \mu\text{m}$$

Cette valeur est si faible qu'elle est impossible à mettre en évidence.

On peut noter qu'*a posteriori* l'hypothèse de considérer la déviation vers l'Est comme une perturbation à la chute libre et la déviation Nord-Sud comme une perturbation à la déviation vers l'Est est parfaitement justifiée par les valeurs numériques obtenues.

A. Application directe du cours

1. Poids d'une valise dans un ascenseur

Un voyageur prend l'ascenseur sa valise à la main. Cette dernière lui semble-t-elle plus ou moins lourde (on suppose que le voyageur la maintient immobile par rapport à l'ascenseur) lorsque l'ascenseur :

1. démarre en montant,
2. démarre en descendant,
3. est arrêté,
4. ralentit en montant,
5. ralentit en descendant.

2. Impesanteur

On considère un satellite en mouvement circulaire uniforme autour de la Terre de masse M_T . On note $\vec{\omega} = \omega_0 \vec{u}_Z$ le vecteur rotation. On appelle $\mathcal{R}_0$ le référentiel (OXYZ) lié à la Terre et $\mathcal{R}$ le référentiel (Sxyz) le référentiel lié au satellite.

On considère un point matériel M de masse m à l'intérieur du satellite et on souhaite étudier son mouvement. On note (x, y, z) ses coordonnées dans $\mathcal{R}$. P est la projection de M sur le plan Sxy et H sa projection sur l'axe OZ.

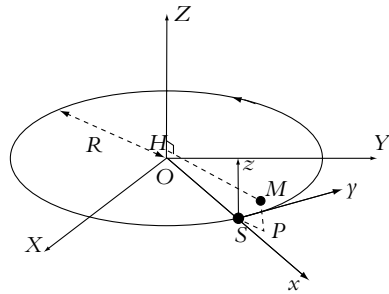


Figure 32.12

1. Montrer que le mouvement de M dans $\mathcal{R}$ est régi par :

$$\begin{cases} \ddot{x} = 3\omega_0^2 x + 2\omega_0 \dot{y} \\ \ddot{y} = -2\omega_0 \dot{x} \\ \ddot{z} = -\omega_0^2 z \end{cases}$$

2. Le point M est lâché sans vitesse initiale en $M_0(x_0, y_0, z_0)$. Résoudre les équations du mouvement et montrer l'existence d'une vitesse de dérive $\vec{v}_D = \langle \vec{v} \rangle$.
3. Proposer deux méthodes permettant de mesurer ω_0 .

3. Forces sur un wagon dans le référentiel terrestre

Un wagon de masse $M = 40.10^3$ kg se déplace à la vitesse $v = 70$ km.h⁻¹ en un lieu de latitude $\lambda = 45^\circ$ Nord sur une voie ferrée horizontale dont la direction fait l'angle θ avec le nord (θ compté positivement vers l'est). On donne $g = 10$ m.s⁻². On considère le référentiel géocentrique $\mathcal{R}_G$ galiléen, et le référentiel terrestre $\mathcal{R}_T$ en rotation (vecteur $\vec{\Omega}$ constant selon l'axe des pôles) par rapport à $\mathcal{R}_G$.

1. Calculer la valeur de Ω .

2. Calculer les composantes de $\vec{\Omega}$ dans un repère (O, X, Y, Z) lié à $\mathcal{R}_T$ avec OX suivant la direction du wagon, OY directement perpendiculaire dans le plan horizontal et OZ suivant la verticale du lieu.
3. Écrire la relation fondamentale de la dynamique dans $\mathcal{R}_T$ en projection dans le repère O, X, Y, Z .
4. Déterminer la direction et le sens de la composante horizontale F_H et de la composante F_V de l'action du wagon sur la voie. Application numérique pour $\theta = \pi/2$.
5. Commenter d'un point de vue pratique les résultats obtenus.

B. Exercices et problèmes

1. Lancement d'une fusée

On étudie le lancement d'une fusée assimilée à un point matériel M de masse m depuis un point P à la surface de la Terre à la latitude λ . On effectue un tir vers le haut avec une vitesse initiale v_0 . On supposera que le référentiel géocentrique est galiléen et que l'accélération de pesanteur peut être considérée suivant la verticale du lieu.

1. Établir les équations du mouvement de la fusée dans le référentiel terrestre. On utilisera la base de projection d'origine P et d'axes Px, Py, Pz tels que Px suivant la direction ouest-est, Py sud-nord et Pz suivant la verticale ascendante.
2. En ne tenant pas compte des effets de la force de Coriolis, déterminer les équations horaires du mouvement.
3. En déduire l'altitude maximale atteinte.
4. En reportant les expressions approchées de la vitesse à l'ordre zéro dans les équations, déterminer les équations horaires approchées du mouvement.
5. Donner la position du point de chute de la fusée.
6. Dans quelle direction a-t-elle été déviée ?
7. Application numérique : on donne $\lambda = 45^\circ$, l'accélération de pesanteur $g = 9,8 \text{ m.s}^{-2}$ et la vitesse initiale $v_0 = 50 \text{ km.h}^{-1}$. Donner les valeurs numériques de la hauteur maximale atteinte ainsi que de la position du point de chute.
8. A quelle vitesse faut-il lancer la fusée pour obtenir une déviation du point de chute de 2,0 m ? Donner la hauteur maximale atteinte dans ce cas.

2. Tir balistique dans le référentiel terrestre

On note $\mathcal{R}_g$ le référentiel géocentrique, supposé galiléen. On note $\mathcal{R}_T$ le référentiel terrestre en rotation par rapport à $\mathcal{R}_g$.

Le vecteur rotation de $\mathcal{R}_T$ par rapport à $\mathcal{R}_g$ est noté $\vec{\Omega}$, suivant l'axe des pôles avec $\Omega = 7,3 \cdot 10^{-5} \text{ rad.s}^{-1}$. En un lieu A de latitude $\lambda = 35^\circ$, un canon tire un obus à la vitesse $v_0 = 1000 \text{ m.s}^{-1}$, vers le nord, avec un angle $\alpha = 45^\circ$ par rapport à l'horizontale.

On désigne par $Axyz$ un repère orthonormé lié à la terre, Ax étant dirigé vers l'est, Ay vers le nord et Az selon la verticale locale. On ne tient pas compte de la résistance de l'air et de la variation de g avec l'altitude. On prendra $g = 9,8 \text{ m.s}^{-2}$. On néglige la variation de latitude au cours du mouvement.

1. Exprimer $\vec{\Omega}$ dans le repère $Axyz$.
2. Écrire la relation fondamentale de la dynamique dans $\mathcal{R}_T$. La projeter sur le repère $Axyz$.
3. Intégrer une fois l'équation en $\ddot{z}$ et celle en $\ddot{y}$. En déduire que l'équation complète selon Ax est :

$$\ddot{x} = -4\Omega^2 x + 2\Omega \cos \lambda g t - 2\nu_0 \Omega \sin \alpha \cos \lambda + 2\nu_0 \Omega \cos \alpha \sin \lambda$$

4. La durée du mouvement est de l'ordre de la minute et la déviation x de l'ordre de la centaine de mètres. Donner une estimation numérique de chacun des termes du membre de droite de l'équation précédente. En déduire une expression approchée de $\ddot{x}$ puis de x à la date t .
5. Déterminer une expression approchée de $\dot{z}$ puis de z à la date t . Déterminer l'instant de retour au sol. Application numérique.
6. Déterminer la déviation de l'obus selon Ax . Effectuer l'application numérique. On entend souvent la remarque selon laquelle un projectile dévie vers la droite dans l'hémisphère nord. Est-ce vrai ? Pour un tir horizontal ($\alpha = 0$), quelle différence y-a-t-il entre les deux hémisphères ?

3. Mouvement dans une station spatiale, d'après ENSTIM 2005

La station spatiale internationale en construction depuis 1998 est située à une altitude d'environ $h = 400 \text{ km}$. Son mouvement est étudié dans le référentiel géocentrique $\mathcal{K}$, d'origine O (centre de la Terre), considéré comme galiléen. Dans $\mathcal{K}$, elle a un mouvement circulaire uniforme de centre O , à la vitesse angulaire ω .

Données : rayon terrestre $R_T = 6400 \text{ km}$; accélération de la pesanteur à la surface du globe $g_0 = 9,8 \text{ m.s}^{-2}$;

1. a) Calculer ω en fonction de g_0 , R_T et h .
- b) Calculer sa période de rotation.

La station spatiale est en rotation synchrone autour de la Terre : elle tourne sur elle-même avec un vecteur vitesse angulaire identique à celui de son mouvement orbital, $\vec{\omega}$.

On désigne par $\mathcal{K}'$, le référentiel lié à la station. L'origine de ce référentiel est situé au centre de masse, S , de la station. L'axe Sx est dirigé suivant $\vec{R} = \vec{OS}$, l'axe Sy est tangent à la trajectoire dans le sens de révolution et l'axe Sz complète le trièdre orthonormé. Dans ce référentiel, un corps ponctuel M , de masse m , est en mouvement dans le plan Sxy . Il est repéré dans la station par le rayon vecteur $\vec{r} = \vec{SM}$.

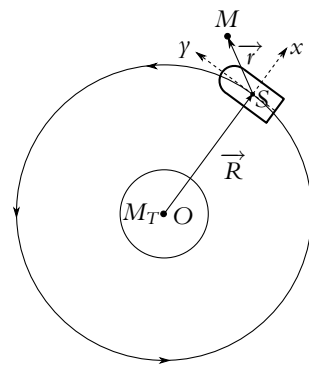


Figure 32.13

2. Pourquoi le référentiel $\mathcal{K}'$ n'est-il pas galiléen ?

3. Définir le point coïncidant à M et donner son accélération $\vec{a}_e(M)$ en fonction de $\vec{r}$, $\vec{R}$ et ω .

En déduire la force d'inertie d'entraînement $\vec{f}_e$ exercée sur la masse m dans $\mathcal{K}'$.

4. Si la particule M est animée d'une vitesse $\vec{v}$ dans $\mathcal{K}'$, quelle force d'inertie supplémentaire lui est appliquée ?

Exprimer cette force en fonction de m , $\vec{\omega}$ et $\vec{v}$.

La particule se trouvant dans le voisinage proche de la station, l'inégalité $r \ll R$ sera toujours vérifiée dans la suite du problème.

5. À l'aide d'un développement limité arrêté au premier ordre en r/R , montrer que la force d'attraction gravitationnelle qu'exerce la Terre sur le corps M s'écrit :

$$\vec{F} = -m\omega^2 \left(\vec{R} + \vec{r} - 3x\vec{u}_x \right)$$

où $\vec{u}_x$ est le vecteur unitaire porté par l'axe Sx et (x, y) sont les coordonnées de $\vec{r}$ dans $\mathcal{K}'$. On rappelle que $(1+x)^\alpha \simeq 1 + \alpha x$ lorsque $|x| \ll 1$.

6. Le corps M est une balle qu'un cosmonaute lance en direction de la Terre avec la vitesse relative $\vec{v}_0 = -v_0\vec{u}_x$ ($v_0 \ll \omega R$) dans $\mathcal{K}'$ depuis l'origine S de ce référentiel.

Établir l'équation du mouvement dans $\mathcal{K}'$ de la balle sous la forme de deux équations différentielles pour les variables x et y .

7. Intégrer ces équations, montrer que la trajectoire suivie est une ellipse et déterminer sa période de parcours.

4. Pendule de Foucault

On s'intéresse au mouvement d'un pendule simple constitué d'une masse $m = 30$ kg suspendue à l'extrémité d'un filin de masse négligeable et de longueur $l = 67$ m. L'autre extrémité du filin est accrochée à un point A fixe par rapport au sol, situé à une hauteur égale à l . À l'instant initial, on écarte le pendule de sa position d'équilibre d'un angle $\alpha = 5^\circ$ dans le plan méridien et on l'abandonne sans vitesse initiale. En un point P de latitude λ , on utilise la base de projection cartésienne $Pxyz$ en prenant Pz selon la direction verticale du lieu et Px dirigé vers l'est.

1. On suppose dans un premier temps que le référentiel terrestre est galiléen.
 - a) Montrer que le mouvement s'effectue dans un plan qu'on précisera.
 - b) Établir l'équation horaire du mouvement par exemple en donnant l'expression de l'angle θ entre le filin et la verticale.
 - c) Calculer les amplitudes maximales des positions, des vitesses et des accélérations dans les deux directions où s'effectue le mouvement.
2. On tient compte dans la suite de la rotation de la Terre sur elle-même. Déterminer la valeur de la vitesse angulaire Ω associée.
3. En déduire que le fait de tenir compte de la rotation de la Terre est une correction par rapport au mouvement précédent.

4. Expliquer qualitativement pourquoi on peut considérer que le mouvement de ce pendule ne détecte pas la rotation de la Terre à l'équateur.
5. Expliciter dans la base de projection proposée les équations du mouvement.
6. En faisant des approximations à justifier à l'aide des questions précédentes, montrer que les équations du mouvement peuvent s'écrire sous la forme :

$$\begin{cases} \ddot{x} - 2\Omega\dot{y} \sin \lambda + \omega_0^2 x = 0 \\ \ddot{y} + 2\Omega\dot{x} \sin \lambda + \omega_0^2 y = 0 \\ T = mg \end{cases}$$

7. On résout ce système en utilisant la notation complexe : on pose $\underline{Z} = x + iy$. En déduire l'équation différentielle vérifiée par $\underline{Z}$.
8. La résoudre pour obtenir les équations horaires $x(t)$ et $y(t)$.
9. Interpréter physiquement la solution.
10. Déterminer la durée d'un tour complet du plan d'oscillations à la latitude $48^\circ 51'$ (latitude de Paris). Cette expérience a été réalisée sous la coupole du Panthéon en 1852 par Léon Foucault (1819-1868) qui mesura une période de 31 h 46 min. Que peut-on penser des résultats obtenus ?
11. Comparer les périodes à l'équateur, aux pôles et à la latitude de 45° .
12. Ce résultat dépend-il de l'hémisphère dans lequel est réalisée l'expérience ?

5. **Modèle statique des marées, d'après Centrale PC 2004**

On considère le système isolé de deux astres en interaction gravitationnelle :

- La Terre (T) de masse $M_T = 6.10^{24}$ kg, de centre T et de rayon $R_T = 6400$ km.
- Un astre attracteur (A) de masse M_A , de centre A : en pratique ce sera la Lune ($M_T = 7,4.10^{22}$ kg) ou le Soleil ($M_T = 2.10^{30}$ kg) :

(A) et (T) sont supposés avoir une distribution de masse sphérique, ce qui assure qu'ils se comportent, pour l'extérieur, comme des masses ponctuelles concentrées en leur centre vis-à-vis des forces de gravitation. On note les différents champs de gravitation créés en un point P extérieur aux astres :

$$\vec{G}_T(P) = -kM_T \frac{\vec{TP}}{TP^3} \quad \text{et} \quad \vec{G}_A(P) = -kM_A \frac{\vec{AP}}{AP^3}$$

avec $k = 6,7.10^{-11} \text{ N.m}^2.\text{kg}^{-2}$.

On définit plusieurs référentiels utiles :

- Le référentiel de Copernic (centré sur le centre de masse du système solaire) noté $\mathcal{R}_C$ et supposé galiléen.
- Le référentiel géocentrique $\mathcal{R}_G$ centré en T , en translation par rapport à $\mathcal{R}_C$.

- Le référentiel $\mathcal{R}_T$ lié au sol terrestre. Son origine sera prise en T . Le mouvement de $\mathcal{R}_T$ par rapport à $\mathcal{R}_G$ est une rotation autour de l'axe des pôles, de vecteur rotation constant $\vec{\Omega} = \Omega \vec{u}_z$.

1. Soit un point P situé à la surface de (T) repéré dans $\mathcal{R}_T$, immobile dans ce référentiel. En plus des forces gravitationnelles et d'inertie, P est soumis à des forces de contact notées $\vec{f}$.

a) Quelle est l'accélération du point T , centre d'inertie de (T) dans $\mathcal{R}_G$?

b) Écrire la relation fondamentale de la dynamique appliquée en P dans $\mathcal{R}_T$.

c) En plus de $\vec{f}$, il apparaît d'autres termes dans l'équation précédente ; on associe ces termes à deux groupes distincts :

- Le premier ne varie pas dans le temps dans $\mathcal{R}_T$ (on rappelle que P est fixe dans $\mathcal{R}_T$), on peut le noter $m\vec{X}(P)$.
- Le deuxième varie dans le temps à cause du mouvement apparent de (A) dans $\mathcal{R}_T$; on le note $m\vec{C}_{AT}(P)$ et on le nomme "force de marée de l'astre (A) en P ".

Proposer une notation plus usitée pour le champ $m\vec{X}(P)$; de quel champ bien connu s'agit-il ? Vérifier que $\vec{C}_{AT}(P) = \vec{G}_A(P) - \vec{G}_A(T)$.

2. Calcul de la force de marée

Soit $\mathcal{P}_A$ le plan dans lequel se déplacent T et A . Pour calculer $\vec{C}_{AT}(P)$ à une date donnée on munit l'espace d'un repère $Txyz$ représenté sur la figure (32.14) (les échelles ne sont pas respectées). Les coordonnées du point P à la surface de la Terre (rayon R_T) sont (x, y, z) , avec $x \ll TA$, $y \ll TA$, $z \ll TA$ et $R_T \ll TA$. La base associée au repère est $(\vec{u}_x, \vec{u}_y, \vec{u}_z)$.

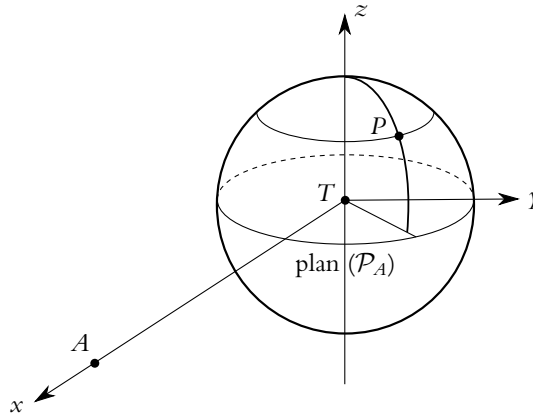


Figure 32.14

a) Calculer $\vec{C}_{AT}(P)$ puis montrer par un développement limité au premier ordre en x/TA , ou y/TA ou z/TA , sur chaque composante du vecteur, qu'on obtient :

$$\vec{C}_{AT}(P) = -\frac{kM_A}{TA^2} \left(\frac{-2x}{TA} \vec{u}_x + \frac{y}{TA} \vec{u}_y + \frac{z}{TA} \vec{u}_z \right)$$

b) Évaluer l'ordre de grandeur de $\|\vec{C}_{AT}\|(P)/\|\vec{G}_T(P)\|$ à la surface de la Terre si (A) est la Lune avec $TL = 3,8 \cdot 10^8$ m.

Représenter sur un schéma clair ce champ $\vec{C}_{AT}(P)$ aux quatre points suivants de $\mathcal{P}_A$: $P_1(R_T, 0, 0)$, $P_2(0, R_T, 0)$, $P_3(-R_T, 0, 0)$ et $P_4(0, -R_T, 0)$.

c) On peut montrer que la force de marée $m\vec{C}_{AT}(P)$ dérive d'une énergie potentielle :

$$E_{p,A}(P) = kmM_A \frac{1}{2TA^3} (-2x^2 + y^2 + z^2)$$

On définit l'angle zénithal de (A) en P par $Z_A(P)$: c'est l'angle (TP, TA) . Montrer que :

$$E_{p,A}(P) = kmM_A \frac{R_T^2}{2TA^3} (1 - 3 \cos^2 Z_A(P))$$

3. Dans l'hypothèse d'une Terre homogène soumise uniquement à la pesanteur, la surface serait sphérique de rayon R_T .

a) Exprimer l'énergie potentielle totale d'un point P à l'altitude z soumis à la pesanteur (qui sera assimilée à $\vec{G}_T(P)$) et à $\vec{C}_{AT}(P)$ en fonction de g accélération de la pesanteur à la surface de la Terre, k , M_A , R_T , TA , $Z_A(P)$, z et m .

b) On fait l'hypothèse que la Terre prend la forme d'une surface équipotentielle, c'est-à-dire que $E_p(P) = E_{p0} = \text{cte}$. Exprimer z pour tout point P et en déduire $z_{\min}$ et $z_{\max}$ sachant que $Z_A(P)$ varie de 0° à 180° . En déduire l'amplitude maximale $\Delta z = z_{\max} - z_{\min}$. Calculer Δz pour la Lune et le Soleil ($TS = 1,5 \cdot 10^{11}$ m et $g = 9,8 \text{ m.s}^{-2}$). Combien de marées hautes peut-on prévoir par jour ?

Thermodynamique

33

Généralités sur les systèmes thermodynamiques

1. Notions de base

1.1 Système thermodynamique

Un système thermodynamique est un système contenant un très grand nombre de particules élémentaires. Ce peut être un fluide (gaz ou liquide) ou un solide ou encore l'ensemble des électrons libres dans un métal...

Il existe 3 niveaux de descriptions : les échelles *macroscopique*, *mésoscopique* et *microscopique*. On retrouve ces mêmes niveaux de descriptions en électrostatique et en magnétostatique par exemple.

- L'échelle *macroscopique* est notre échelle. À cette échelle, tout milieu paraît continu.
- L'échelle *microscopique* est celle des particules élémentaires du système. À cette échelle, la matière est discontinue.
- L'échelle *mésoscopique* est intermédiaire. Un petit volume $d\tau$ à cette échelle (par exemple 1 mm^3), est assez grand pour contenir un très grand nombre de particules (environ 10^{17} pour l'air sous 1 bar à 0°C) et qu'on puisse considérer le milieu comme continu et assez petit pour considérer que les grandeurs macroscopiques (pression, densité, température) y sont uniformes. Pour satisfaire ces conditions, le volume $d\tau$ doit être choisi grand devant le volume moyen occupé par la particule au cours de son mouvement.

1.2 Les différentes descriptions

a) Description microscopique

Son utilisation relève de la thermodynamique statistique. Le système est constitué de N particules (typiquement de l'ordre du nombre d'Avogadro, $N_A = 6,02 \cdot 10^{23} \text{ mol}^{-1}$). Si on veut traiter le problème grâce à la mécanique classique, en supposant les particules ponctuelles, il faut étudier pour chaque particule sa position (3 variables d'espace) et sa vitesse (3 autres variables), soit aux total $6N$ variables, tout en connaissant les interactions entre particules : ce problème n'est pas soluble. Si on ne peut considérer les particules ponctuelles, le nombre de variables augmente encore. On ne peut donc pas suivre la trajectoire d'une particule au cours du temps, le traitement ne peut être que *statistique*. Dans le cas d'un gaz sous faible pression, les particules suivent un mouvement de type *mouvement brownien* : la trajectoire semble constituée de segments de droite successifs dont les directions sont aléatoires. On peut aisément observer ce type de trajectoire avec de la fumée dans le faisceau lumineux d'une lampe. Il existe des méthodes de simulation sur ordinateur avec quelques milliers de particules qui donnent de bons résultats.

b) Description macroscopique

Elle décrit un comportement collectif et résulte de l'observation des phénomènes : elle est *phénoménologique*. La moyenne des effets microscopiques donne à toute grandeur un aspect continu. Cependant, le fait d'avoir affaire à un grand nombre de particules impose d'introduire de nouvelles grandeurs par rapport à la mécanique comme la *température* (notée T) et l'*entropie* (notée S) qu'on étudiera dans le chapitre sur le deuxième principe. L'introduction de ces grandeurs est le résultat du traitement statistique sous-jacent évoqué précédemment ; on développera partiellement ces notions dans le chapitre sur l'entropie statistique (Cf. chapitre 39).

On peut imaginer le système de simulation suivant décrit par la figure 33.1 : un cylindre contenant des billes d'acier, fermé à sa partie inférieure par un piston mobile, et à sa partie supérieure par une membrane permettant d'enregistrer les chocs des billes venant la percuter. Le piston a un mouvement oscillatoire.

S'il y a peu de billes ou si le piston va lentement, peu de billes viendront percuter la membrane pendant une seconde et on pourra enregistrer chaque choc indépendamment (Cf. figure 33.2) ; en revanche, s'il y a beaucoup de billes et que l'oscillation est d'amplitude et de fréquence suffisantes, on ne pourra distinguer chaque choc car un nombre

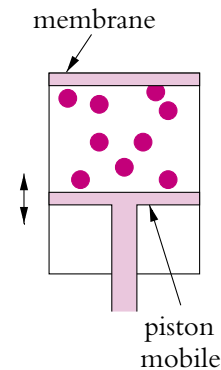


Figure 33.1 Simulation : billes dans un cylindre.

important de billes viendra percuter la membrane par seconde, on ne pourra mesurer que la force moyenne s'exerçant sur la paroi (Cf. figure 33.3).

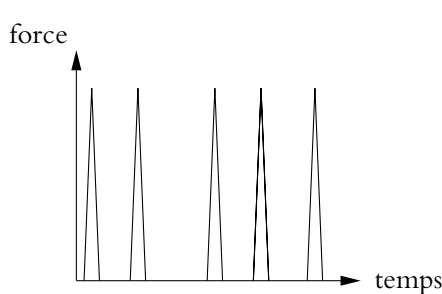


Figure 33.2 Force exercée dans le cas d'une seule bille.

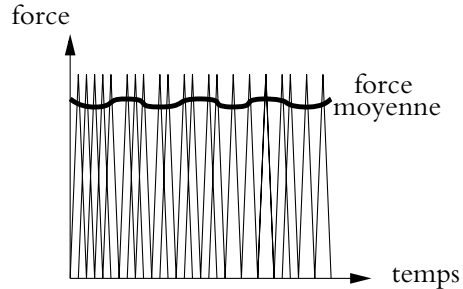


Figure 33.3 Force exercée dans le cas de nombreuses billes.

1.3 Équilibre thermodynamique

L'hypothèse d'*équilibre thermodynamique* est la suivante : il n'y a **pas de mouvement macroscopique** à l'intérieur du système thermodynamique considéré mais il peut y avoir un **mouvement d'ensemble du système**.

Attention, il y a un mouvement microscopique : les particules sont en mouvement incessant, c'est ce qu'on appelle l'*agitation thermique*.

Alors les grandeurs macroscopiques (température, pression...) du système sont indépendantes du temps. Si on y regarde de plus près, les grandeurs macroscopiques ne sont pas tout à fait constantes, elles fluctuent autour d'une valeur moyenne, comme le montre l'exemple précédent avec les billes si l'appareil est assez sensible pour enregistrer les différents chocs (Cf. figure 33.3). La force moyenne n'est pas rigoureusement constante car pendant un intervalle de temps dt , il n'y pas toujours exactement le même nombre de particules qui vient heurter la membrane.

En fait, la notion d'équilibre thermodynamique **n'est pas indépendante de la durée d'observation**. Par exemple, une motte de beurre qu'on laisse à l'extérieur d'un réfrigérateur semblera peu évoluer pendant quelque temps et on pourra considérer qu'elle est à l'équilibre. Cependant, au bout d'un temps important, elle se sera effondrée.

Un système peut être considéré à l'équilibre pendant un certain temps (ses variables macroscopiques n'évoluent pas de manière visible), alors que si on l'observe sur un temps beaucoup plus long, il évolue.

Ce problème d'évolution du système sera en partie traité dans le cours de seconde année.

1.4 Densité d'un système

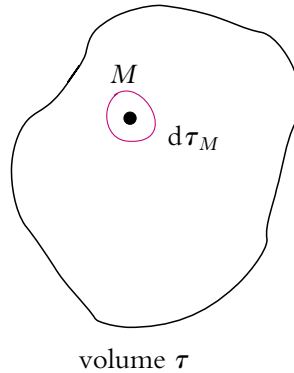


Figure 33.4 Volume mésoscopique $d\tau_M$ autour du point M .

Soit un volume τ contenant N particules et un point M de ce volume. On s'intéresse à un volume mésoscopique $d\tau_M$ autour du point M . Soit dN_M le nombre de particules dans ce volume $d\tau_M$. On définit la densité particulaire $n(M)$, c'est-à-dire le nombre de particules par unité de volume par :

$$dN_M = n(M) d\tau_M \quad \text{ou} \quad n(M) = \frac{dN_M}{d\tau_M} \quad (33.1)$$

La densité particulaire $n(M)$ s'exprime donc en m^{-3} .

1.5 Pression

On reviendra plus en détail sur la pression en étudiant la statique des fluides. Il s'agit juste ici de donner une définition.

Soit un fluide en équilibre thermodynamique et en contact avec une paroi solide.

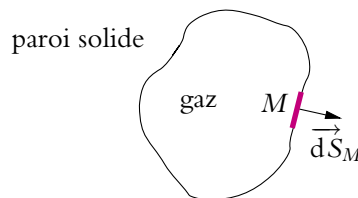


Figure 33.5 Surface élémentaire, vecteur surface.

Soit dS_M une surface élémentaire de cette paroi solide (Cf. figure 33.5) et $d\vec{S}_M = dS_M \vec{n}_M$ le vecteur surface associé, $\vec{n}_M$ étant normal à la surface en M . Les conventions d'orientation de $d\vec{S}_M$ sont détaillées dans le cours d'électromagnétisme. On signale seulement que pour une surface fermée, le vecteur est orienté vers l'extérieur. La composante normale de la force exercée par le gaz sur cette surface

$\vec{dF}_n(M)$ est proportionnelle à la norme dS_M de $\vec{dS}_M$. On définit la pression $P(M)$ régnant au point M par :

$$\vec{dF}_n(M) = P(M)\vec{dS}_M \quad (33.2)$$

La pression est une grandeur positive et indépendante de l'orientation de $\vec{dS}_M$.

On supposera dans la suite qu'il existe une pression en tout point du fluide même si ce point n'est pas en contact avec une paroi (c'est la pression que mesurerait un manomètre placé en ce point).

L'unité de pression est le pascal (Pa) : $1 \text{ Pa} = 1 \text{ N.m}^{-2}$. C'est une petite unité, aussi emploie-t-on souvent le bar : $1 \text{ bar} = 10^5 \text{ Pa}$. On utilise aussi le millimètre de mercure (en référence au baromètre à mercure), la pression moyenne au niveau de la mer étant $1,013 \text{ bar} = 760 \text{ mmHg}$.

2. Température

La notion physiologique de température est connue de chacun : elle correspond à une sensation de « chaud » ou de « froid » (termes qui s'avèreront incorrects par la suite car ils correspondent à des échanges d'énergie et non à une caractéristique intrinsèque). Ainsi, un barreau de bois paraît plus chaud qu'un barreau de fer même si leur température est identique. En fait, le fer conduit mieux la chaleur que le bois et, quand on touche du fer plus froid que le corps, ce dernier fournit de l'énergie au barreau qui est plus vite évacuée que dans le cas du bois.

On va préciser ici la définition de la température et la rattacher à un phénomène physique de façon à en faire une grandeur mesurable.

2.1 Équilibre thermique

On place un tube capillaire (tube très fin) rempli de liquide coloré dans un bain-marie : on observe alors que le niveau du liquide s'élève (par phénomène de dilatation) jusqu'à se stabiliser à une certaine hauteur. Quand l'évolution cesse, on dit qu'il y a *équilibre thermique* entre le bain-marie et le capillaire. Tous les paramètres décrivant l'état du système gardent alors une valeur constante.

On peut alors énoncer le principe zéro de la thermodynamique :

Tous les systèmes en équilibre thermique avec un système donné sont en équilibre thermique entre eux.

On peut le vérifier grâce à un ensemble de tubes capillaires identiques, en mettant en contact tous les systèmes, chacun avec son capillaire : les niveaux respectifs ne varient pas : on dit qu'**ils ont la même température**. Ceci donne une définition relative de la température : on sait reconnaître les systèmes qui ont la même température et ceux qui ont des températures différentes.

2.2 Repérage de la température

Intuitivement, il y a une relation directe entre le niveau du liquide dans un tube capillaire donné et la sensation de « chaud » correspondante. La hauteur L du liquide est un paramètre qui varie en fonction de la température : il s'agit d'une *grandeur thermométrique*. Le capillaire constitue alors un thermomètre, permettant de *repérer* la température. Plus L est grand, plus la sensation de « chaud » est intense. On note θ la valeur numérique évaluant la notion de température : à chaque valeur de L , on peut faire correspondre une valeur de θ , par l'intermédiaire d'une fonction croissante f :

$$L \xrightarrow{f} \theta$$

On peut, par exemple, choisir pour $f : L \xrightarrow{f} \theta = aL^2 + bL + c$, avec $a, b > 0$. θ est appelée *température empirique*, $\theta = f(L)$ est l'*équation thermométrique* qui la définit.

L'avantage à ce stade est que l'on peut maintenant faire référence à θ sans savoir quel capillaire a été utilisé. Encore faut-il disposer de « points fixes », appelés aussi *repères thermométriques*, c'est-à-dire de systèmes qui gardent toujours la même température, pour pouvoir étalonner chaque thermomètre du type capillaire même si le diamètre varie. On reviendra sur ce problème des repères thermométriques dans le paragraphe 2.6.

2.3 Échelles centésimales linéaires

Soit L la hauteur de liquide dans un capillaire. Deux conditions doivent être remplies pour parler d'échelle centésimale linéaire :

- la fonction f est choisie linéaire : $\theta = \alpha L + \beta$,
- les repères thermométriques sont choisis de manière bien déterminée :
 1. mélange eau liquide + eau solide sous la pression atmosphérique :
 $\theta = \theta_0 = 0$ (θ n'a pas d'unité, c'est un repérage),
 2. mélange eau liquide + eau vapeur sous la pression atmosphérique :
 $\theta = \theta_{100} = 100$.

On utilise la propriété suivante (voir chapitre 40) : à pression fixée, un corps pur en équilibre sous deux états différents à pression donnée a toujours la même température. On remarquera qu'on a besoin de deux valeurs arbitrairement fixées, l'échelle sera dite à *échelle à deux points fixes*. Ces deux conventions permettent ainsi de déterminer les coefficients α et β de l'équation thermométrique. Dans le cas d'un capillaire, cela correspond, du point de vue pratique, à repérer deux hauteurs L_0 et L_{100} puis à tracer entre elles 99 graduations équidistantes.

Pour un thermomètre quelconque caractérisé par une grandeur thermométrique G (résistance, longueur d'un fil...) :

$$\theta = \alpha G + \beta$$

avec $\theta_0 = \alpha G_0 + \beta$ et $\theta_{100} = \alpha G_{100} + \beta$,
on en tire :

$$\theta = 100 \frac{G - G_0}{G_{100} - G_0} \quad (33.3)$$

On obtient directement l'équation thermométrique (33.3) à l'aide de la pente :

$$\frac{\theta - 0}{G - G_0} = \frac{100 - 0}{G_{100} - G_0}$$

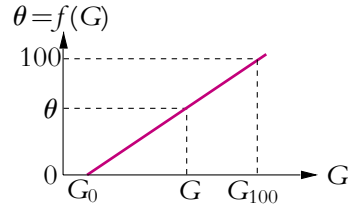


Figure 33.6 Échelle centésimale linéaire.



- Ce repérage ne repose sur aucune base théorique. La nature même de ce qu'on appelle température n'est pas élucidée (c'est comme si on définissait l'intensité du courant électrique comme la lecture de l'ampèremètre sans dire que $\delta q = I dt$)
- Rien ne prouve que le comportement du thermomètre sera effectivement linéaire en fonction du phénomène physique responsable de la sensation de « chaud » ou « froid ».
- On ne peut *a priori* faire le rapport de deux températures : il ne fait pas deux fois plus chaud à 40 degrés qu'à 20 degrés.

Pour définir la température à partir d'une base physique, il faut étudier les propriétés des gaz aux très faibles pressions.

2.4 Propriétés des gaz aux faibles pressions

Il a été mis en évidence que pour des pressions très faibles, de l'ordre du pascal, tous les gaz ont des propriétés simples et identiques.

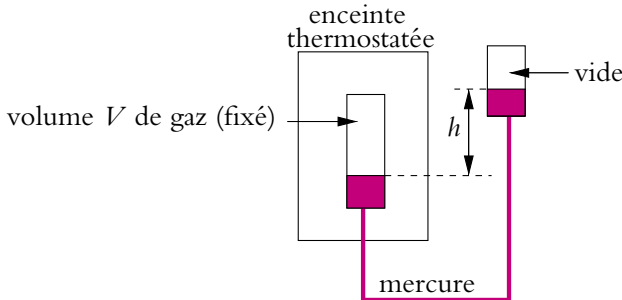


Figure 33.7 Thermomètre à gaz.

Soit une masse m donnée de gaz, maintenue à température constante dans une enceinte thermostatée.

On étudie l'évolution de la pression et du volume de ce gaz grâce au dispositif de la figure 33.7. On verra dans le chapitre 34 que la pression du gaz est proportionnelle à la hauteur h de la colonne de mercure entre les deux récipients avec : $P = \mu gh$, μ étant la masse volumique de mercure et g l'accélération de la pesanteur.

Ainsi, si on étudie l'évolution du produit PV de plusieurs gaz avec le dispositif de la figure 33.7 pour une température θ , on peut extrapoler les isothermes $PV = f(P)$ pour $P \rightarrow 0$ (pressions « évanescentes »), on obtient une limite finie $(PV)_\theta$. Cette valeur limite, pour une mole de gaz, ne dépend pas de la nature du gaz (Cf. figure 33.8). Sur la figure, les traits pleins correspondent aux mesures et les pointillés aux extrapolations.

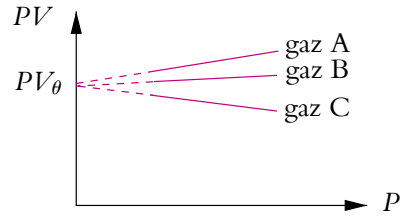


Figure 33.8 Extrapolation des isothermes pour une mole de gaz.

On dispose donc d'une grandeur physique ($\lim_{P \rightarrow 0} PV$ pour 1 mole) qui ne dépend que de la température, quel que soit le gaz considéré.

2.5 Température absolue

a) Principe

On applique les résultats du paragraphe précédent. Ainsi, aux pressions évanescentes, pour une masse m fixée (quelconque) d'un gaz donné, le produit PV ne dépend que de la température et est proportionnel au nombre de molécules donc, pour deux températures θ_1 et θ_2 différentes, on peut écrire :

$$\frac{(PV)_{\theta_1}}{(PV)_{\theta_2}} = f(\theta_1, \theta_2)$$

Le rapport ne dépend que des températures. Par exemple, on mesure que ce rapport est égal à 1,3661 pour $\theta_1 = 100$ et $\theta_2 = 0$ quel que soit le gaz et quelle que soit la masse de ce gaz.

Par définition, ce rapport est égal au rapport des températures absolues T_1 et T_2 correspondant aux températures empiriques θ_1 et θ_2 .

Il suffit donc de se donner par convention la valeur de la température pour un repère thermométrique pour en déduire toutes les autres valeurs. Si T_R est cette température de référence, on peut écrire :

$$T = T_R \frac{(PV)_\theta}{(PV)_{\theta_R}} \quad (33.4)$$

La définition de T ne nécessitant qu'une seule valeur arbitraire, l'échelle est dite *échelle à un point fixe*.

b) Échelle légale

Elle est définie par le point fixe utilisé, qui est **le point triple de l'eau** (quand coexistent les trois phases eau liquide, eau vapeur et glace, la température est fixée, de même que la pression comme on le verra dans le chapitre sur les changements d'état). On pose :

$$T_R = 273,16 \text{ K (kelvin)} \quad (33.5)$$

Pourquoi la valeur de 273,16 K? Tout simplement parce que ce choix permet de trouver $T_0 = 273,15 \text{ K}$ pour $\theta = 0$ et $T_{100} = 373,15 \text{ K}$ pour $\theta = 100$, soit un écart de 100 unités comme dans les échelles centésimales. La température absolue est une grandeur mesurable, directement proportionnelle à une grandeur physique (le produit PV aux pressions évanouissantes) et en combinant les équations (33.4) et (33.5), on obtient la relation :

$$T = 273,16 \frac{(PV)_T}{(PV)_{273,16}} \quad (33.6)$$

► **Remarque** : De même qu'on ne dit pas « degré kelvin » mais « kelvin », on n'écrit pas °K mais K.

c) Échelle Celsius

Cette échelle résulte d'un décalage de 273,15 unités à partir de l'échelle légale :

$$t(^{\circ}\text{C}) = T(\text{K}) - 273,15$$

Elle coïncide donc avec les échelles centésimales pour les deux points fixes $t = 0^{\circ}\text{C}$ et $t = 100^{\circ}\text{C}$ (mais ce sont *a priori* les seules températures de même valeur dans les deux types d'échelles).

2.6 Mesure des températures

a) Thermomètre légal

C'est un thermomètre à gaz (H_2) comme celui de la figure 33.7 qui sert à déterminer la température absolue d'un certain nombre de points fixes (mélange de 2 ou 3 phases d'un même corps pur). On pourra ensuite étalonner à l'aide de ces points fixes d'autres thermomètres plus maniables.

b) Thermomètres usuels

L'ensemble de points fixes ainsi déterminés constitue l'*Echelle Internationale de Température (EIT)*. Le thermomètre à dihydrogène n'est pas d'un emploi très aisé, et, depuis 1927, une échelle légale de températures a donc été élaborée à partir de différentes points fixes. On utilise actuellement l'*Echelle Internationale de Température de 1990 (EIT-1990)*. Voici quelques-uns de ses points fixes :

Température (en K) du point fixe	Corps	Type de point
24,5561	néon	Point triple
273,16	eau	Point triple
505,078	étain	Point de congélation
1357,77	cuivre	Point de congélation

Tableau 33.1 Différents points fixes d'après EIT-90.

En particulier, pour les températures usuelles, l'EIT-90 définit :

- entre 13,803 K et 1 234,93 K, 14 points fixes et l'instrument d'interpolation est un thermomètre à résistance de platine ;
- au-dessus de 1 234,93 K, 3 points fixes et la température est mesurée par pyrométrie optique en utilisant les lois de rayonnement des corps (ces lois seront étudiées en deuxième année MP).

De même, dans la pratique, on utilise aussi divers types de thermomètres suivant le domaine de température considéré : thermomètres à dilatation à mercure ou alcool, thermomètres à résistance de platine et thermocouples (industrie), thermomètres à gaz, hélium ou dihydrogène, (très basses températures), pyromètres optiques (très hautes températures). Ces thermomètres sont étalonnés à partir des points fixes.

La grandeur thermométrique n'est en général pas une fonction linéaire de la température et on utilise le plus souvent une interpolation par un polynôme de degré 3 au moins pour représenter correctement l'équation thermométrique.

3. Équation d'état

3.1 Variables d'état

Les variables d'état sont des grandeurs macroscopiques permettant de définir l'état du système à un instant donné, comme par exemple la pression P , le volume V , la température T ou la charge q .

Il existe deux types de variables d'état : si on considère un système homogène et que l'on isole un volume V par la pensée, certaines variables sont proportionnelles à V , d'autres non. On appelle :

- *variable extensive*, une variable proportionnelle à V ; par exemple, le volume V , la charge q , l'énergie E , la masse m sont des variables extensives ;
- *variable intensive*, une variable indépendante de V ; par exemple, la pression P , la température T , la densité d , la masse volumique μ , le potentiel U sont des variables intensives.

Les variables extensives sont additives : ainsi si deux systèmes A et B ont pour énergie E_A (respectivement E_B) et température T_A (respectivement T_B), l'énergie du système $A \cup B$ sera $E = E_A + E_B$ car E est une variable extensive mais, par contre, la température du système $A \cup B$ ne sera pas $T_A + T_B$.

On appelle *variables indépendantes*, les variables nécessaires et suffisantes à la description d'un système.

Par exemple, pour un gaz, il suffit de connaître le nombre de moles n , la température T et le volume V pour déterminer l'état du système (pas besoin de connaître la pression qui peut être déterminée grâce aux autres variables).

On peut alors redéfinir l'équilibre thermodynamique, en écrivant qu'un système est à **l'équilibre si ses variables d'état sont constantes au cours du temps**.

3.2 Équation d'état

Ce qui est intéressant, c'est de prévoir les valeurs prises par certaines variables d'un système, connaissant les autres par la mesure, par exemple. Il est donc nécessaire de rechercher des relations entre les variables d'états. C'est ce qui a été fait pour les gaz par les physiciens Robert Boyle (1627-1691) et Edme Mariotte (vers 1620-1684) qui se sont intéressés à la l'évolution du produit PV en fonction de la température, puis par Jacques Charles (1746-1823) qui a étudié la variation de la pression en fonction de la température à volume fixé et Louis-Joseph Gay-Lussac (1778-1850) qui a étudié la variation du volume en fonction de la température à pression fixée.

On peut prendre un exemple autre qu'un gaz : on désire connaître la longueur L d'un fil élastique, sachant que sa température est T et qu'il est soumis à une force F . C'est possible car il existe une relation reliant L , F et T telle que : $f(F, T, L) = 0$. Cela signifie que les trois variables F , T et L ne sont pas indépendantes.

On appelle *équation d'état* la relation $f(F, T, L) = 0$.

Pour un fluide homogène non chargé, tel qu'un gaz, l'équation d'état est de la forme $f(n, P, V, T) = 0$.

3.3 Gaz parfait

Un gaz parfait est un gaz dont les variables P (Pression), V (Volume), T (Température) et n (nombre de moles) sont reliées par l'équation d'état :

$$PV = nRT \quad (33.7)$$

R est la constante des gaz parfaits : $R = 8,314 \text{ J.K}^{-1}.\text{mol}^{-1}$.

Il est important de noter que dans l'équation (33.7), les différentes grandeurs sont exprimées dans les unités du Système International :

- P en pascal, • V en mètre cube,
- T en kelvin, • n en mole.

On peut exprimer l'équation des gaz parfaits (33.7) avec la constante de Boltzmann $k_B = 1,38 \cdot 10^{-23} \text{ J.K}^{-1}$ en utilisant les relations suivantes :

$$n = \frac{N}{N_A} \quad \text{et} \quad k_B = \frac{R}{N_A}$$

où N est le nombre de particules du gaz et $N_A = 6,02 \cdot 10^{23} \text{ mol}^{-1}$ le nombre d'Avogadro. Alors l'équation (33.7) devient :

$$PV = Nk_B T \quad (33.8)$$

On remarque que l'on retrouve les lois décrites dans le paragraphe 2.4 pour tous les gaz aux pressions évanouissantes. On en déduit que tous les gaz se comportent comme des gaz parfaits aux faibles pressions.

En fait, on verra dans la suite que le modèle du gaz parfait est une approximation supposant qu'il n'y a aucune interaction entre les particules du gaz, celles-ci étant considérées comme ponctuelles. Cependant, on utilise ce modèle car, en général, il donne des résultats proches de la réalité lorsque la pression n'est pas trop importante.

3.4 Écart au gaz parfait ; gaz réels

Comme cela vient d'être signalé, le gaz parfait modélise un gaz idéal sans interactions entre les particules. En fait, le physicien Van der Waals (vers 1873) a proposé une autre équation d'état qui porte son nom. Il apparaît effectivement dans certaines conditions nécessaire de tenir compte des interactions entre particules, ces interactions portent d'ailleurs le nom de forces de Van der Waals.

Comment peut-on se rendre compte du comportement non parfait d'un gaz ?

Le plus simple est de tracer des isothermes dans le diagramme d'Amagat, où on représente la pression P en fonction du produit PV (Cf. figure 33.9) :

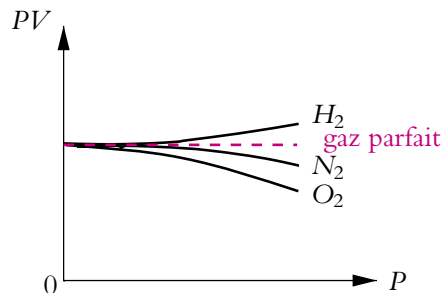


Figure 33.9 Écart du comportement des gaz réels par rapport aux gaz parfaits.

D'après l'équation du gaz parfait, si la température T est constante, le produit PV l'est aussi : on obtient donc des droites horizontales. Cela est vérifié aux faibles pressions, mais lorsque la pression augmente, on s'aperçoit que les isothermes ne sont plus des droites : le gaz ne se comporte alors plus comme un gaz parfait. Van der Waals a donc proposé l'équation suivante :

$$\left(P + \frac{n^2 a}{V^2}\right) (V - nb) = nRT \quad (33.9)$$

où a et b sont deux constantes positives, caractéristiques du gaz, dont on va chercher la signification. Les termes correctifs $\frac{n^2 a}{V^2}$ et nb sont, en général, très faibles devant respectivement P et V . Si on les néglige, on retrouve l'équation du gaz parfait.

La constante b est homogène à un volume molaire et est appelée *covolume*. Si on réécrit l'équation (33.9) en négligeant le terme contenant a :

$$V - nb = \frac{nRT}{P}$$

Ainsi, lorsqu'on augmente la pression très fortement (à la limite $P \rightarrow \infty$), on voit que le terme $V - nb$ tend vers 0 et que le volume tend vers nb . Si on modélise les particules comme des boules (modèle dit des sphères dures), lorsque la pression est très grande, les particules se touchent : **nb représente donc le volume minimal occupé par les particules.**

Le terme $\frac{n^2 a}{V^2}$ qui est homogène à une pression porte le nom de *pression moléculaire*. Il est lié à l'existence d'une force attractive qui s'exerce entre les particules (la force de Van der Waals). En effet, si on se place à une pression courante pour laquelle on peut négliger nb devant V , on peut exprimer la pression par :

$$P = \frac{nRT}{V} - \frac{n^2 a}{V^2}$$

On remarque que le premier terme du membre $\frac{nRT}{V}$ correspond à la pression d'un gaz parfait. La pression du gaz est donc inférieure à celle d'un gaz parfait de même volume et même température. Or c'est cette même pression que le gaz exerce sur l'enceinte qui l'entoure ; le gaz a ainsi moins tendance à « pousser » l'enceinte, cela signifie qu'il y a une force attractive entre les particules qui n'existe pas dans les gaz parfaits.

L'équation (33.9) rend bien compte du comportement des gaz réels jusque dans l'état liquide à l'exception du dihydrogène et de l'hélium qui ont un comportement différent des autres gaz comme le montre la figure 33.9. Les molécules de ces gaz sont très petites, elles peuvent donc s'approcher très près les unes des autres et les interactions

deviennent alors répulsives. Il existe d'autres équations d'état qui s'appliquent dans certaines conditions mais qu'on ne donne pas ici.

3.5 Coefficients thermoélastiques

Ce sont les coefficients accessibles à l'expérience et qui permettent d'établir l'équation d'état. En effet, le comportement d'un gaz constitué de n moles dépend des trois variables P , T et V . Pour étudier expérimentalement son comportement, il est plus aisé de fixer l'une des trois variables, et de s'intéresser à l'influence de l'une des deux autres sur la troisième. Quand on maintient une variable constante, celle-ci est notée en indice dans une dérivée partielle.

Dans le cadre du programme, on s'intéresse à deux coefficients thermoélastiques :

- le *coefficient de compressibilité isotherme* : $\chi_T = -\frac{1}{V} \left(\frac{\partial V}{\partial P} \right)_T$,
- le *coefficient de dilatation isobare* : $\alpha = \frac{1}{V} \left(\frac{\partial V}{\partial T} \right)_P$,

Lors de petites variations de P , T ou V , on peut assimiler les dérivées aux rapports des variations. Ainsi lors d'une petite augmentation de pression dP à température constante, la variation relative de volume est : $\frac{dV}{V} = -\chi_T dP$. De même, lors d'une faible augmentation de température dT à pression constante, la variation relative de volume est : $\frac{dV}{V} = \alpha dT$.

On peut remarquer que α est de signe quelconque alors que χ_T est toujours positif. On verra en effet dans le cours de seconde année PC-PSI qu'une des conditions de stabilité de l'équilibre thermodynamique est $\left(\frac{\partial V}{\partial P} \right)_T < 0$, ce qui explique la présence du signe $-$ dans l'expression.

► **Remarque** : Il suffit de connaître ces deux coefficients pour déterminer l'équation d'état.

Pour terminer, on peut exprimer ces deux coefficients dans le cas du gaz parfait en utilisant l'équation (33.7) :

$$\begin{aligned} \left(\frac{\partial V}{\partial T} \right)_P &= \left(\frac{\partial}{\partial T} \left(\frac{nRT}{P} \right) \right)_P = \frac{nR}{P} \quad \text{donc} \quad \alpha = \frac{1}{V} \frac{nR}{P} = \frac{1}{T} \\ \left(\frac{\partial V}{\partial P} \right)_T &= \left(\frac{\partial}{\partial P} \left(\frac{nRT}{P} \right) \right)_T = \frac{-nRT}{P^2} \quad \text{donc} \quad \chi_T = \frac{1}{V} \frac{nRT}{P^2} = \frac{1}{P} \end{aligned}$$

A. Applications directes du cours

1. Pression de pneumatiques

En hiver, pour une température extérieure de -10°C , un automobiliste règle la pression de ses pneus à 2 atm, pression préconisée par le constructeur.

1. Quelle mesure donne un manomètre ?
2. Quelle serait son indication en été à 30°C ? On suppose que le volume des pneus ne varie pas et qu'il n'y a aucune fuite au niveau de ce dernier.
3. Calculer la variation relative de pression due au changement de température. Conclure.

2. Thermomètre à mercure

Un thermomètre à mercure est destiné à être utilisé entre 0°C et 150°C . La dilatabilité moyenne du mercure entre 0°C et $t^{\circ}\text{C}$ est donnée par $\frac{V - V_0}{V_0 t} = a - bt + ct^2$ où a , b et c sont des constantes positives.

1. Définir l'échelle affine centésimale $V = A + B\theta$ en exprimant θ en fonction de a , b , c et t .
2. Exprimer la quantité $\Delta = \theta - t$.
3. Montrer que Δ passe par un extremum pour deux températures t_1 et t_2 sachant que $\theta = t$ à 150°C .

3. Échelle Fahrenheit

L'échelle Fahrenheit de température en usage dans plusieurs pays anglo-saxons se déduit de l'échelle Celsius par une transformation affine. Par définition, on a $32^{\circ}\text{F} = 0^{\circ}\text{C}$ et $212^{\circ}\text{F} = 100^{\circ}\text{C}$.

1. Établir la loi permettant le passage des degrés Fahrenheit aux degrés Celsius.
2. Établir la loi réciproque.
3. Convertir 451°F en $^{\circ}\text{C}$.
4. A quelle température, les deux échelles donnent-elles la même indication ?

4. Équation réduite de Van der Waals

On considère un gaz de Van der Waals dont l'équation d'état est :

$$P = \frac{RT}{V - b} - \frac{a}{V^2}$$

On définit le point critique comme le point pour lequel on a :

$$\left(\frac{\partial P}{\partial V}\right)_T = \left(\frac{\partial^2 P}{\partial V^2}\right)_T = 0$$

1. Etablir les équations liant les pression P_C , température T_C et volume V_C critiques à a , b et R .

2. On pose $\Pi = \frac{P}{P_C}$, $\theta = \frac{T}{T_C}$ et $\varphi = \frac{V}{V_C}$. Montrer qu'on a

$$\left(\Pi + \frac{3}{\varphi^2}\right)(3\varphi - 1) = 8\theta$$

3. Quel est l'intérêt d'une telle expression ?

B. Exercices et problèmes

1. Équilibre d'un piston

Un cylindre vertical fermé aux deux bouts est séparé en deux compartiments égaux par un piston homogène, se déplaçant sans frottement. La masse du piston par unité de surface est $\sigma = 1360 \text{ kg.m}^{-2}$. Les deux compartiments contiennent un gaz parfait à la température $t_1 = 0^\circ\text{C}$. La pression qui règne dans le compartiment supérieur est égale à $P_1 = 10 \text{ cmHg}$. L'intensité de la pesanteur est $g = 10 \text{ m.s}^{-2}$.

1. En écrivant que le piston est à l'équilibre, déterminer la pression, en bars, du gaz dans le compartiment du bas.

2. On porte les deux compartiments à $t_2 = 100^\circ\text{C}$. De combien se déplace le piston ?

2. Étude d'un compresseur

Un compresseur est constitué de la façon suivante : un piston se déplace dans un cylindre C qui communique par des soupapes s et s' respectivement avec l'atmosphère (pression P_a) et avec le réservoir R contenant l'air comprimé. Le réservoir R contient initialement de l'air considéré comme gaz parfait à la pression $P_0 \geq P_a$.

Le volume du réservoir R , compte-tenu des canalisations, est V . Le volume offert au gaz dans C varie entre un volume maximum V_M et un volume minimum V_m , volume nuisible résultant de la nécessité d'allouer un certain espace à la soupape s .

- La soupape s s'ouvre lorsque la pression atmosphérique P_a devient supérieure à la pression dans le cylindre C et se ferme pendant la descente du piston.
- La soupape s' s'ouvre lorsque la pression dans le réservoir devient supérieure à celle du gaz dans le cylindre C et se ferme pendant la montée du piston.

Au départ, le piston est dans sa position la plus haute ($V = V_M$), s' est fermée, s est ouverte et le volume V_m est rempli d'air à la pression P_a .

1. a) En supposant que le piston se déplace assez lentement pour que l'air reste à température constante, calculer le volume V'_1 pour lequel s' s'ouvre en fonction de P_0 , P_a et V_M .

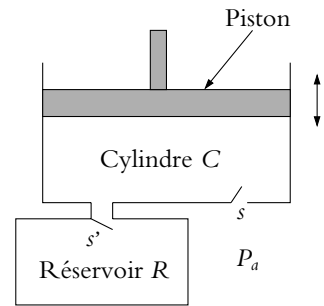


Figure 33.10

- b) Calculer la pression P_1 dans le réservoir R après le premier aller et retour.
- c) En écrivant une condition sur V_1' , calculer la valeur P_{max} au-dessus de laquelle la pression ne peut pas monter dans le réservoir.
2. Calculer la pression P_n dans le réservoir R après n allers et retours du piston.
3. Donner la valeur limite de P_n quand $n \gg 1$. Comparer cette limite avec P_{max} .
4. Calculer P_1 et P_{max} avec $V = 5 \text{ L}$, $V_M = 0.25 \text{ L}$, $V_m = 10 \text{ cm}^3$, $P_0 = P_a = 1 \text{ bar}$.

3. Énergie interne d'un gaz - Gaz réel ou parfait ?

Les valeurs expérimentales de l'énergie interne massique de la vapeur d'eau sont les suivantes :

$T \text{ (K)}$	523	573	623	673
à $P = 10 \text{ bars}$	2711	2793	2874	2956
à $P = 20 \text{ bars}$	2683	2773	2859	2944

1. Tracer les courbes donnant l'énergie interne en fonction de la température.
2. A-t-on un gaz parfait ? Justifier.
3. Comparer les capacités thermiques à celle d'un gaz parfait monoatomique.

Statique des fluides dans le champ de pesanteur

34

1. Définition

D'après le dictionnaire « Le Robert », un fluide est un « corps qui épouse la forme de son contenant ». Parmi les systèmes thermodynamiques définis dans le chapitre précédent, il s'agit des gaz et des liquides. On s'intéresse dans ce chapitre à l'équilibre d'un fluide dans le champ de pesanteur uniforme terrestre. À l'échelle mésoscopique, le volume $d\tau_Q$ autour du point Q , défini dans le chapitre précédent, porte le nom de *particule de fluide*. Il est petit à l'échelle macroscopique mais contient suffisamment de molécules pour qu'on puisse considérer le milieu continu et définir des grandeurs moyennées sur $d\tau_Q$.

Dans le référentiel terrestre $\mathcal{R}$, supposé galiléen, un fluide de masse volumique ρ et limité par un volume τ (Cf. figure 34.1) est soumis à deux types de forces dans le champ de pesanteur :

- une force répartie en volume : le volume mésoscopique $d\tau_Q$ entourant le point Q est soumis au champ de pesanteur uniforme, il subit la force $d\vec{F}_{\text{ext}}(Q)$:

$$d\vec{F}_{\text{ext}}(Q) = dm_Q \vec{g} \quad (34.1)$$

où dm_Q est la masse du volume $d\tau_Q$; le volume τ , de masse m , est alors soumis à la force :

$$\vec{F}_{\text{ext}} = \iiint_{Q \in \tau} dm_Q \vec{g} = m \vec{g}$$

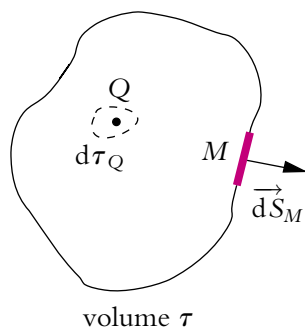


Figure 34.1 Volume et surface élémentaires.

- une force répartie en surface : l'ensemble du volume τ est soumis à une force de pression surfacique exercée par le milieu extérieur ; sur la surface élémentaire entourant le point M, la force exercée sur le fluide est

$$d\vec{F}(M) = -P(M)d\vec{S}_M$$

Le volume τ , entouré de la surface $\mathcal{S}$, est alors soumis à la force :

$$\vec{F}_p = \oint_{M \in \mathcal{S}} -P(M)d\vec{S}_M$$

Comme on l'a indiqué dans le paragraphe précédent, la pression s'exerce en n'importe quel point du fluide et le module de la force élémentaire correspondante en un point donné est le même quelle que soit la direction.

2. Pression d'un fluide soumis au champ de pesanteur

2.1 Relation fondamentale de la statique des fluides

On considère un fluide de masse volumique ρ , en équilibre mécanique dans le référentiel $\mathcal{R}$ terrestre supposé galiléen. Comme cela vient d'être signalé dans le paragraphe précédent, il est soumis, en plus des forces de pression internes au fluide, à son poids.

Soit un volume mésoscopique $d\tau_M$ de forme parallélépipédique, autour d'un point M, de surface horizontale S , compris entre les plans de cote z et $z + dz$ donc de volume $d\tau_M = Sdz$, et de masse volumique $\rho(M)$. On étudie l'équilibre de ce volume dans le référentiel $\mathcal{R}$.

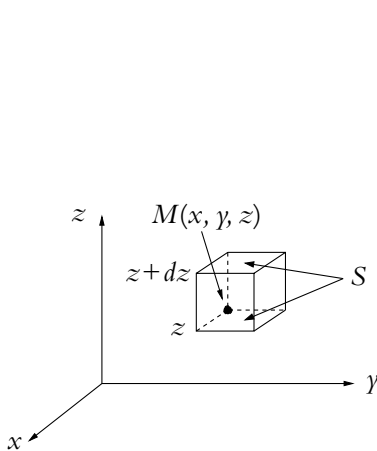


Figure 34.2 Volume mésoscopique $d\tau_M$.

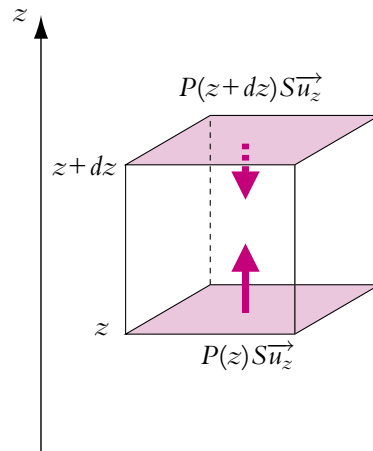


Figure 34.3 Sens des forces de pression.

Les forces s'exerçant sur $d\tau_M$ sont :

- les forces de pression exercées par le reste du fluide sur la surface de $d\tau_M$, notées $d\vec{F}_p$,
- le poids, $d\vec{F}_{\text{ext}} = \rho(M)d\tau_M \vec{g}$.

La relation fondamentale de la dynamique donne à l'équilibre :

$$d\vec{F}_{\text{ext}} + d\vec{F}_p = \vec{0} \quad (34.2)$$

On projette la relation (34.2) sur l'axe Oz . La figure 34.3 indique le sens des projections des forces de pression. Aucune force extérieure n'a de composante selon Ox et Oy , la pression ne dépend donc que de z .

La composante suivant Oz de la résultante des forces de pression, qui s'exercent sur les surfaces S horizontales (l'une à la cote z et l'autre à la cote $z + dz$), a pour expression :

$$dF_{pz} = P(z)S - P(z + dz)S$$

La quantité dz a été choisie petite devant la distance caractéristique de variation de la pression, on peut donc effectuer un développement limité de l'expression précédente au premier ordre en dz :

$$P(z) - P(z + dz) = -\frac{dP}{dz} dz$$

Soit :

$$dF_{pz} = (P(z) - P(z + dz)) S = -\frac{dP}{dz} dz S$$

La relation (34.2) projetée sur Oz donne alors :

$$-\frac{dP}{dz} dz S - \rho g S dz = 0 \Rightarrow -\frac{dP}{dz} - \rho g = 0$$

On obtient l'équation fondamentale de la statique des fluides dans le cas d'un fluide en équilibre dans le champ de pesanteur uniforme :

$$\frac{dP}{dz} = -\rho g \quad (34.3)$$

2.2 Fluide incompressible et homogène

On appelle *fluide incompressible* un fluide dont la masse volumique ρ est indépendante de la pression c'est-à-dire dont le coefficient de compressibilité isotherme χ_T est nul.

On appelle *fluide homogène* un fluide dont la masse volumique ρ est la même en tout point de l'espace. Les liquides sont des bons exemples de tels fluides ; c'est une approximation dont on verra qu'elle est justifiée (voir le calcul en fin de paragraphe).

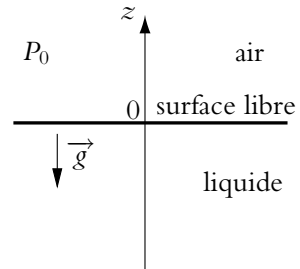


Figure 34.4 Surface libre d'un fluide incompressible.

On considère un liquide incompressible et homogène soumis au champ de pesanteur $\vec{g}$ uniforme. Le liquide est en contact avec l'atmosphère de pression P_0 au niveau de sa surface libre, située en $z = 0$.

Si on intègre l'équation (34.3), sachant que ρ et g sont constants, on obtient :

$$P = -\rho g z + C \quad \text{où } C \text{ est une constante} \quad (34.4)$$

Pour déterminer la constante, on utilise la continuité de la pression au niveau de la surface libre : $P(z = 0) = P_0$. Alors $C = P_0$ et

$$P(z) = P_0 - \rho g z \quad (34.5)$$

Si on désire déterminer la pression à une profondeur h , il suffit d'appliquer l'équation (34.5) pour $z = -h$. On obtient alors l'expression appelée *équation barométrique* :

$$P = \rho g h + P_0 \quad (34.6)$$

Que représente physiquement le terme $\rho g h$ dans l'équation (34.6) ? C'est le poids de la colonne de fluide de hauteur h et de surface 1 m^2 . Ainsi, lorsqu'on se trouve sous une hauteur h d'eau, on ressent sur ses épaules le poids de la colonne d'eau au-dessus de soi en plus de la force de pression exercée par l'atmosphère.

Que doit-on retenir de l'équation barométrique (34.6) ?

1. À une profondeur h , on ressent le poids de la colonne de fluide au-dessus de soi.
2. La pression dans un fluide incompressible homogène croît avec la profondeur, proportionnellement à la masse volumique du fluide.
3. Les surfaces isobares sont d'altitude constante.

On peut calculer la profondeur (sous l'eau) à laquelle on ressent une pression double de la pression atmosphérique. On cherche donc la profondeur h à laquelle $P = 2 P_0$ avec pour l'eau $\rho = 10^3 \text{ kg.m}^{-3}$ et $g \simeq 10 \text{ m.s}^{-2}$. L'équation barométrique (34.6) donne :

$$h = \frac{P_0}{\rho g} \simeq \frac{10^5}{10^3 \times 10} = 10 \text{ m} \quad (34.7)$$

Ainsi sous 10 m d'eau environ, la pression est doublée.

En conclusion, on ne peut pas considérer la pression constante dans un fluide, comme l'eau ou le mercure par exemple. Ceci est dû en particulier à la masse volumique importante du fluide.

2.3 Validité du modèle incompressible

Jusqu'ici on a considéré que le fluide était incompressible. Si l'on se place dans le cas de l'eau, est-ce que cette approximation est justifiée ? La variation de masse volumique est reliée à la variation de volume d'une masse m donnée de fluide. On utilise le

coefficient de compressibilité isotherme défini dans le chapitre précédent qui relie la variation de pression à la variation de volume, soit pour l'eau :

$$\chi_T = -\frac{1}{V} \left(\frac{\partial V}{\partial P} \right)_T = 5.10^{-10} \text{ Pa}^{-1}$$

Si on assimile les petites variations de P et V aux dérivées, la variation relative de volume d'une masse m donnée d'eau est :

$$\frac{\Delta V}{V} = -\chi_T \Delta P \quad (34.8)$$

D'autre part, on sait que la masse m est égale au produit $m = \rho V$ pour un fluide homogène donc, pour une masse m donnée constante, si on différencie logarithmiquement cette expression, on trouve :

$$d(\ln m) = d(\ln \rho) + d(\ln V) = 0 \Leftrightarrow \frac{d\rho}{\rho} + \frac{dV}{V} = 0 \quad (34.9)$$

De plus, l'équation barométrique donne une variation ΔP égale à $\rho g \Delta z$ pour une variation Δz de la profondeur. Ainsi, si on combine ce résultat avec les équations (34.8) et (34.9), pour une variation Δz de la profondeur, on obtient une variation relative de masse volumique :

$$\frac{\Delta \rho}{\rho} = \chi_T \rho g \Delta z \quad (34.10)$$

Dans le cas de l'eau, l'expression (34.10) donne, pour $\Delta z = 1 \text{ m}$:

$$\frac{\Delta \rho}{\rho} = 5.10^{-10} \times 10^3 \times 10 = 5.10^{-6} = 5.10^{-4}\%$$

Ainsi pour avoir une variation relative de 5 % de la masse volumique de l'eau, il faut s'enfoncer à 10 km sous la surface libre : **on peut donc négliger la variation de la masse volumique de l'eau.**

2.4 Applications de la statique des fluides incompressibles

Ce sont des applications de la loi barométrique (34.6). La plupart des applications se déduisent des propriétés des *vases communicants* (tube en U), en particulier les siphons et les baromètres (Cf. figures 34.5 et 34.6).

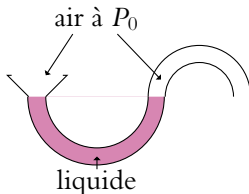


Figure 34.5 Schéma d'un siphon d'évacuation.

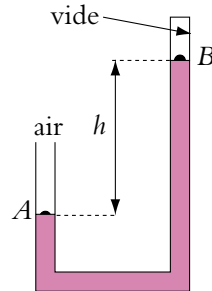


Figure 34.6 Baromètre de Torricelli.

a) Siphon

Le siphon, que l'on trouve sur les évacuations d'eau, évite la remontée des odeurs grâce à de l'eau qui reste en permanence dans le tube coudé. La pression au-dessus de chaque surface libre étant la pression atmosphérique P_0 , **la dénivellation entre les deux surfaces libres est nulle**. On peut toujours utiliser ce principe pour siphonner un récipient, c'est-à-dire vider un récipient plein dans un autre, situé plus bas, à l'aide d'un tuyau (Cf. figure 34.7). Il suffit d'amorcer le syphon en aspirant un peu de liquide, puis le liquide s'écoulera du récipient A vers le récipient B de manière à ce que la dénivellation entre les deux surfaces libres puisse s'annuler, ce qui n'arrivera pas avant que A ne soit vide.

b) Baromètre

Le baromètre à mercure le plus courant est le baromètre de Torricelli (Cf. figure 34.6). Il s'agit d'un tube en U fermé à une extrémité, l'autre extrémité étant ouverte à l'air libre (pression P_0). Du côté de l'extrémité fermée, le mercure est surmonté de vide (pression nulle). Si on applique la relation barométrique entre les points A et B, on a : $P_A = P_B + \rho gh$. Or $P_A = P_0$ et $P_B = 0$ donc $P_0 = \rho gh$. On peut donc lire directement la pression atmosphérique en lisant la hauteur de la dénivellation de mercure, d'où l'unité de pression en centimètres de mercure. Pourquoi n'utilise-t-on pas d'eau dans un baromètre ? Au paragraphe précédent, on a vu qu'il fallait 10 m d'eau pour avoir une différence de pression égale à 1 bar, il faudrait donc un baromètre d'au moins 10 m de haut.

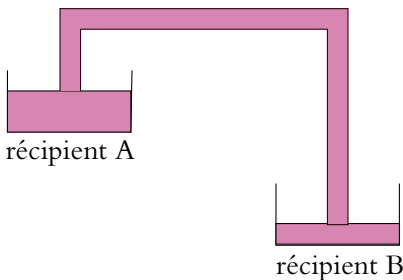


Figure 34.7 Principe du siphonnage.

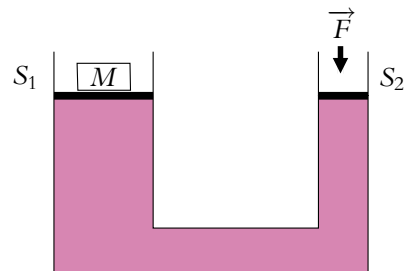


Figure 34.8 Schéma d'un vérin.

c) Vérin hydraulique

La dernière application présentée est très utile pour le levage d'objets lourds, il s'agit du *vérin hydraulique*. Le schéma de principe est représenté sur la figure 34.8. Il s'agit d'un tube dont les colonnes de droite et de gauche ont des sections différentes, respectivement S_1 et S_2 ($S_1 > S_2$). À l'équilibre, seule la pression atmosphérique P_0 s'exerce et les surfaces libres sont au même niveau. On pose une masse M à gauche.

Quelle force supplémentaire F doit-on exercer pour maintenir les surfaces libres toujours au même niveau ? On applique la loi barométrique et, puisque la dénivellation h est nulle, les pressions P_A et P_B de chaque côté sont identiques. Si on néglige la masse des pistons on a :

$$P_A = P_0 + \frac{Mg}{S_1} \quad \text{et} \quad P_B = P_0 + \frac{F}{S_2}$$

soit

$$F = \frac{S_2}{S_1} Mg < Mg \quad (34.11)$$

Ainsi, la force à exercer pour maintenir la masse M est inférieure à Mg et la force nécessaire à la soulever sera elle aussi inférieure à ce qu'elle serait sans le vérin.

2.5 Gaz parfait : atmosphère isotherme

a) Détermination de la loi $P(z)$ pour l'atmosphère isotherme

On s'intéresse maintenant au cas où **le fluide est un gaz parfait compressible**. On va calculer la pression dans un gaz parfait soumis à son poids comme seule force extérieure. On fait l'hypothèse que le gaz est en équilibre isotherme, c'est-à-dire que sa **température T est constante**.

On rappelle l'équation (34.3) traduisant l'équilibre d'un volume mésoscopique (Cf. paragraphe 2) :

$$\frac{dP}{dz} = -\rho g \quad (34.12)$$



Attention, contrairement au cas du liquide incompressible, on ne peut intégrer cette équation directement car ρ n'est pas constante ici. Il faut utiliser l'équation des gaz parfaits, y faire apparaître ρ puis exprimer ρ en fonction de P .

Pour un volume homogène V , de masse M_V , par définition la masse volumique : $\rho = \frac{M_V}{V}$. Or si on appelle N le nombre de particules et m la masse d'une particule, $M_V = Nm$, d'où la transformation de l'équation du gaz parfait :

$$P = \frac{Nk_B T}{V} \Rightarrow P = \frac{M_V k_B T}{mV} = \frac{\rho k_B T}{m} \quad (34.13)$$

L'expression précédente peut aussi être écrite en fonction de la masse molaire M :

$$P = \frac{\rho R T}{M}$$

En reportant l'expression (34.13) dans l'équation (34.12) on obtient :

$$\frac{dP}{dz} = -\frac{mg}{k_B T} P \Leftrightarrow \frac{dP}{P} = -\frac{mg}{k_B T} dz \quad (34.14)$$

En appelant P_0 la pression à l'altitude $z = 0$, on intègre l'équation précédente entre $z = 0$ et z pour trouver :

$$P = P_0 \exp\left(-\frac{mgz}{k_B T}\right) \quad (34.15)$$



Attention, on peut intégrer directement (34.14) uniquement parce que T est considérée constante. Si on prend un modèle plus réaliste de l'atmosphère pour lequel T dépend de z , il faudra remplacer T par son expression en fonction de z avant d'intégrer l'équation (34.14).

On définit la *hauteur d'échelle* H par l'altitude à laquelle P_0 est divisée par e (base du logarithme népérien). Alors :

$$P = P_0 \exp\left(-\frac{z}{H}\right) \quad \text{avec} \quad H = \frac{k_B T}{mg} \quad (34.16)$$

On peut aussi écrire : $H = \frac{RT}{Mg}$.

Dans le cas de l'air à température ordinaire ($M = 29 \text{ g.mol}^{-1}$ et $T = 273 \text{ K}$), on trouve $H \simeq 8 \text{ km}$. De même, la pression à une hauteur de 10 m est $P = P_0 \exp\left(-\frac{10}{8} \times 10^3\right) = 0,99875 P_0$; la variation relative de la pression sur une hauteur de 10 m est inférieure à $0,2 \%$.

Ainsi, sur des hauteurs faibles devant H , hauteur caractéristique de variation de la pression en fonction de l'altitude, on peut négliger la variation de pression dans un gaz avec l'altitude. C'est le contraire de ce qui se passe pour un liquide.

Cela est dû à la faible masse volumique des gaz. Dans un gaz, on supposera donc la pression uniforme à l'équilibre tant qu'on l'étudie sur des hauteurs faibles devant H .

b) Interprétation de la loi $P(z)$ pour l'atmosphère isotherme.

En utilisant la loi des gaz parfaits et la relation de proportionnalité entre le nombre de particules par unité de volume n et ρ , la relation (34.15) devient :

$$n(z) = n_0 \exp\left(-\frac{mgz}{k_B T}\right)$$

Ainsi :

- n décroît avec l'altitude ;
- n décroît d'autant plus rapidement que les particules sont lourdes (m plus grand) ;
- n décroît d'autant plus rapidement que T est petit (on observe en effet que l'épaisseur de l'atmosphère varie entre le jour et la nuit).

Dans l'exponentielle apparaît le rapport de deux termes homogènes à une énergie :

- mgz , l'énergie potentielle d'une particule dans le champ de pesanteur g ,
- $k_B T$, l'énergie d'agitation thermique comme ce sera vu dans le chapitre suivant.

Ainsi la répartition des particules avec l'altitude résulte de la compétition entre l'agitation thermique qui tend à uniformiser la répartition et le poids qui tend à attirer les particules vers le bas.

Cette loi se généralise. Si on considère un système de n_0 particules par unité de volume, en équilibre thermique à la température T , le nombre de ces particules ayant une énergie potentielle E_p est :

$$n = A \exp \left(-\frac{E_p}{k_B T} \right)$$

où A est une constante. Le terme exponentiel porte le nom de *facteur de Boltzmann*. Ainsi, statistiquement, les zones de plus basse énergie sont les plus peuplées. La loi d'Arrhénius vue en cinétique chimique peut s'interpréter comme une application de ce résultat.

2.6 Conclusion

En conclusion de ces paragraphes, il faut retenir que la pression varie fortement dans un liquide incompressible mais que sa variation avec l'altitude est très faible dans un gaz.

3. Actions exercées par les fluides au repos. Poussée d'Archimède

3.1 Théorème d'Archimède

La tradition rapporte que le physicien grec Archimède de Syracuse (−287, −212 av. J.-C.) aurait découvert (en prononçant son célèbre *ευρηκα* (*eureka*), qui signifie « j'ai trouvé ») en prenant un bain pourquoi les objets flottent. Ce qui fut autrefois un principe se démontre grâce à la loi fondamentale de la statique des fluides. On étudie un objet de forme quelconque, de volume V , de masse M , immergé dans un fluide au repos. On a vu précédemment avec la loi de la statique des fluides que l'expression de la pression ne dépend que du champ de pesanteur $\vec{g}$. Dans les deux cas des figures 34.9 et 34.10, la présence ou non de l'objet ne changeant pas le champ de pesanteur, le champ de pression est lui aussi inchangé.

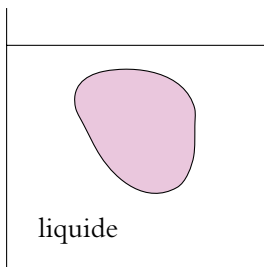


Figure 34.9 Corps solide de volume V plongé dans un liquide.

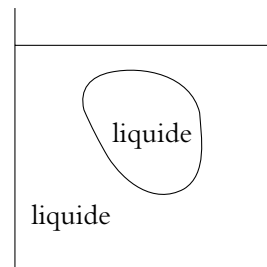


Figure 34.10 Liquide occupant le même volume V .

Ainsi, la force de pression s'exerçant sur le volume V de la figure 34.10 due au reste du fluide est la même que celle exercée sur l'objet occupant le même volume V de la figure 34.9. Si on fait un bilan semblable à celui du paragraphe 2 mais sur le volume V total de la figure 34.10, ce volume est en équilibre sous l'action de deux forces :

- la résultante des forces de pression exercées sur la surface de fluide $\vec{F}_p$.
- le poids de ce volume V de fluide, qui est l'unique force extérieure.

La relation fondamentale de la dynamique appliquée au volume V de masse m_f de fluide donne donc :

$$\vec{F}_p + m_f \vec{g} = \vec{0}$$

soit

$$\vec{F}_p = -m_f \vec{g}$$

D'après de qui précède, puisque la répartition de pression est inchangée dans le fluide, la force $\vec{\Pi}_A$ exercée sur le solide de la figure 34.9 est égale à $\vec{F}_p$ soit :

$$\vec{\Pi}_A = -m_f \vec{g} \quad (34.17)$$

où m_f est la masse de fluide du volume V de la figure 34.10, qui est aussi la masse de fluide déplacé par l'introduction de l'objet dans la figure 34.9, et donc $m_f \vec{g}$ est le poids du volume de fluide déplacé. Cette force porte le nom de *poussée d'Archimède*, c'est la résultante des forces de pression s'exerçant sur le solide.

On peut énoncer le théorème d'Archimède :

Tout corps immergé au repos subit de la part du fluide une force opposée à celle du poids du volume de fluide déplacé.

Ainsi sur un corps de masse m , qui est en équilibre dans un fluide (par exemple dans l'eau), deux forces s'exercent :

- son poids $m \vec{g}$ vers le bas, s'exerçant au centre de gravité de l'objet,
- la poussée d'Archimède $-m_f \vec{g}$ vers le haut, s'exerçant au centre de gravité du volume de fluide déplacé.

On rappelle que la poussée d'Archimède est la résultante des forces de pression sur le solide étudié, de surface $\mathcal{S}$:

$$\vec{\Pi}_A = \oint_{M \in \mathcal{S}} -P(M) d\vec{S}_M = -m_f \vec{g}$$

Le centre d'inertie du volume de fluide déplacé C , point d'application de la poussée d'Archimède, est aussi appelé *centre de poussée*. Dans le cas de la figure 34.11, le centre de poussée est au-dessous du centre de gravité G . Le couple de force fait pivoter l'objet qui se retourne (Cf. figure 34.12). Le centre de poussée se trouve alors au-dessus du centre de gravité. Le couple fait alors osciller le solide autour de sa position d'équilibre stable où G et C sont sur la même verticale.

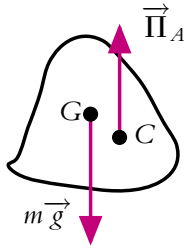


Figure 34.11 C au-dessous de G . Le solide bascule.

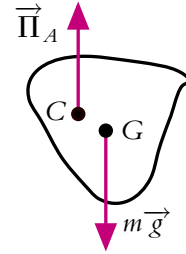


Figure 34.12 C au-dessus de G . Le solide oscille autour de sa position d'équilibre.



- C'est la variation de pression avec l'altitude qui est à l'origine de la poussée d'Archimède, la pression étant plus importante sous l'objet qu'au-dessus.
- Un corps plus léger que le fluide dans lequel il est immergé flotte puisque son poids est inférieur à celui du fluide déplacé.
- Quand un corps est en mouvement lent dans un fluide, on peut considérer que le champ de pression est pratiquement le même qu'à l'équilibre et que la force de pression exercée par le fluide sur l'objet est égale à la force à l'équilibre, c'est-à-dire à la poussée d'Archimède.

► **Remarque :** Si l'objet est en équilibre entre deux fluides, comme il y a continuité de la pression à l'interface entre les deux fluides, le résultat obtenu pour la poussée d'Archimède se généralise. Il suffit dans ce cas d'ajouter les poussées d'Archimède dues aux deux fluides.

3.2 Quelques applications de la poussée d'Archimède

a) L'iceberg

On veut calculer le volume émergé d'un iceberg, supposé homogène, de masse m . On appelle :

- V_i le volume immergé,
- V_e le volume émergé,
- ρ_l la masse volumique de l'eau liquide,
- ρ_g la masse volumique de la glace.

Les forces exercées sur l'iceberg sont :

1. son poids $m\vec{g} = \rho_g(V_i + V_e)\vec{g}$,
2. la poussée d'Archimède due à l'eau $-\rho_l V_i \vec{g}$,
3. la poussée d'Archimède due à l'air $-\rho_{\text{air}} V_e \vec{g}$.

L'équation d'équilibre mécanique s'exprime donc en projection sur la verticale :

$$mg - \rho_l V_i g - \rho_{\text{air}} V_e g = 0$$

Sachant que la masse volumique de l'air est très inférieure à celle de l'eau (environ 1000 fois), on peut négliger la poussée d'Archimède exercée par l'air devant celle exercée par l'eau. L'équation précédente devient alors après simplification :

$$\rho_g(V_i + V_e) = \rho_l V_i \Rightarrow \frac{V_e}{V_i} = \frac{\rho_l}{\rho_g} - 1 = \frac{1}{0.9} - 1 \simeq \frac{1}{9} \quad (34.18)$$

Ainsi, seulement un dixième du volume d'un iceberg est émergé.

b) Le glaçon

Il s'agit d'un problème semblable à celui de l'iceberg, mais posé différemment. On prend un verre d'eau, dans lequel on plonge un glaçon de masse m . On remplit ensuite le verre d'eau à ras bord. Est-ce que le verre d'eau risque de déborder à cause de la fonte du glaçon ?

La question à se poser est : est-ce que le volume d'eau déplacé par la partie immergée du glaçon est supérieure ou inférieure à celui du glaçon une fois fondu ? Lorsque le glaçon est fondu, sa masse volumique est celle de l'eau liquide ρ_l et le volume occupé est $V = \frac{m}{\rho_l}$. Si on écrit l'équation d'équilibre du glaçon, on obtient :

$$mg - \rho_l V_i g = 0$$

Le volume de la partie immergée est donc $V_i = \frac{m}{\rho_l}$ soit le même que celui occupé par le glaçon une fois fondu. **Le niveau d'eau ne bouge donc pas.**

c) Un ballon ascensionnel

On étudie le cas simple d'un ballon dont le volume V est constant, gonflé avec un gaz plus léger que l'air atmosphérique (hélium ou air chaud) de masse volumique ρ_{gaz} . On appelle m la masse du ballon (enveloppe, gaz et nacelle) et on néglige le volume de la nacelle devant celui du ballon. En considérant le référentiel terrestre galiléen, les forces qui s'exercent sur le ballon sont la poussée d'Archimède $\vec{\Pi}_A = -\rho_{\text{air}} V \vec{g} = -m_{\text{air}} \vec{g}$ et son poids $m\vec{g}$. La relation fondamentale de la dynamique appliquée au ballon donne :

$$m\vec{a} = \vec{\Pi}_A + m\vec{g} \quad (34.19)$$

En projection sur l'axe Oz orienté vers le haut, elle s'écrit :

$$m \frac{dz}{dt} = (m_{\text{air}} - m)g$$

Ainsi, si le ballon est plus léger que l'air déplacé, $m_{\text{air}} - m$ est positif et le ballon subit une accélération verticale ascendante. Il faut noter que, dans ce cas, on admet, comme cela a été signalé dans le paragraphe précédent, que la force subie par l'objet est la même à l'équilibre et en mouvement.

3.3 Cas où on ne peut pas appliquer la poussée d'Archimède

Lorsqu'une partie de la surface du corps immergé n'est pas en contact avec le fluide (Cf. figure 34.13), il est inexact d'écrire que la résultante des forces de pression est égale à la poussée d'Archimède. Il faut, dans ce cas, effectuer le calcul de la force de pression en revenant à la définition et en intégrant la force sur la surface du corps.

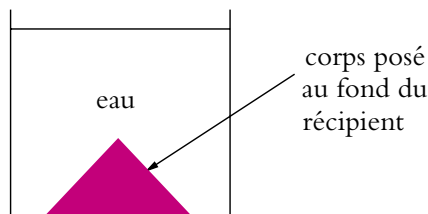


Figure 34.13 Cas de non application du théorème d'Archimède.

De même, quand un solide obstrue le fond d'un récipient troué et rempli d'un liquide, la résultante des forces de pression sur le bouchon n'est pas égale à la poussée d'Archimède : le bouchon n'est en équilibre ni dans le fluide, ni dans le fluide et l'air puisqu'il est en contact avec le récipient, ce qui introduit une différence de pression de part et d'autre du fond de celui-ci. Il faut, dans ce cas encore, effectuer le calcul de la force de pression en revenant à la définition et en intégrant la force sur la surface du corps.

A. Applications directes du cours

1. Interface huile-eau dans un tube en U

On considère un tube en U rempli d'eau jusqu'à une hauteur $H = 10$ cm par rapport au fond. La section du tube est $s = 1$ cm². On ajoute 2 cm³ d'huile dans une des branches du tube. La masse volumique de l'huile est égale à 0.6 fois celle de l'eau.

A quelle hauteur se trouve l'interface entre l'eau et l'huile? A quelle hauteur se trouve la surface libre de l'eau dans l'autre branche?

2. Remontée d'un plongeur

1. Avant une remontée rapide vers la surface, les plongeurs sous-marins voient leur poumons de l'air qu'il contient. Pourquoi?

2. En supposant que l'air contenu dans les poumons est à la température du corps (37 °C) et que son volume est de 3 L à 10 m de profondeur, calculer son volume à la surface. La pression à la surface est $P_0 = 1$ bar, la masse volumique de l'eau est $\rho = 10^3$ kg.m⁻³. On prendra $g = 9,8$ m.s⁻².

3. Pression au sommet de l'Everest

On considère que la température de l'air (gaz parfait) décroît linéairement avec l'altitude. Au niveau de la mer la température vaut 20 °C, et au sommet de l'Everest (altitude 8850 m), elle vaut -40 °C. La masse molaire de l'air est $M = 29$ g.mol⁻¹.

- Déterminer la loi de variation de la température avec l'altitude.
- Établir la loi de variation de la pression avec l'altitude. En déduire la pression au sommet de l'Everest en fonction de la pression au niveau de la mer.
- Montrer que dans l'atmosphère, on a une relation du type $PV^k = \text{constante}$. Calculer k .

4. Force de pression sur un barrage

On s'intéresse à un barrage constitué d'un mur droit vertical. La hauteur d'eau est $h = 5$ m. La largeur du cours d'eau est $\ell = 4$ m. Calculer la force exercée par l'eau sur le barrage sachant que la pression de l'air est $P_0 = 1$ bar. La masse volumique de l'eau est $\rho = 10^3$ kg.m⁻³. On prendra $g = 9,8$ m.s⁻².

5. Pression dans une fosse océanique

Dans tout l'exercice, on suppose $g = 9,8$ m.s⁻² constant.

1. Quelle est la pression dans une fosse océanique à 10 km de profondeur, en supposant que l'eau de mer est incompressible et de masse volumique $1,03 \cdot 10^3$ kg.m⁻³ et que la pression à la surface est $P_0 = 1$ bar?

2. On considère maintenant que l'eau a une masse volumique à la surface $\rho_0 = 1,03 \cdot 10^3$ kg.m⁻³ et un coefficient de compressibilité isotherme $\chi_T = 4,5 \cdot 10^{-5}$ bar⁻¹.

a) Exprimer χ_T en fonction de la masse volumique ρ .

- b) À partir de l'équation fondamentale de la statique, établir l'équation différentielle vérifiée par p . Déterminer p en fonction de z .
3. a) En déduire l'équation différentielle donnant la pression et la résoudre.
- b) Calculer la pression à une profondeur de 10 km et comparer avec le cas précédent.

B. Exercices et problèmes

1. Tube tournant

Un tube coudé plonge dans de l'eau (masse volumique $\mu = 10^3 \text{ kg.m}^{-3}$). Ce tube tourne autour de l'axe vertical Oz à la vitesse angulaire ω . La pression atmosphérique est notée P_a . On note ρ la masse volumique de l'air supposée constante, la section du tube est supposée être très faible.

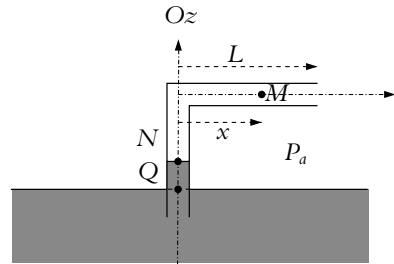


Figure 34.14

- On raisonne dans le référentiel lié au tube, qui n'est pas galiléen. L'air est en équilibre dans ce référentiel. On considère une petite tranche d'air de masse dm , comprise entre x et $x + dx$.
 - Quelles sont les forces qui s'exercent sur cette tranche ?
 - On fait l'hypothèse que l'effet du poids est négligeable. En s'inspirant de la démonstration de la loi fondamentale de la statique, établir une équation différentielle reliant la dérivée de la pression $\frac{dp}{dx}$, ρ , ω et x .
 - Calculer la pression $P(x)$ de l'air dans le tube à une distance x de l'axe Oz en fonction de P_a , ω , L et x .
- Déterminer la dénivellation du liquide.

2. Fluide dans un cylindre en rotation

Un flacon cylindrique de hauteur 40 cm et de diamètre 10 cm contient de l'eau sur une hauteur de 10 cm quand le flacon est au repos. On le fait tourner autour de son axe à vitesse angulaire ω constante.

- Quelle est l'équation caractéristique de la surface libre ?
- Pour quelle valeur de ω le centre du cylindre est-il découvert ?
- Pour quelle valeur de ω l'eau ne déborde-t-elle pas ?
- Quelle est la pression au fond du flacon si $\omega = 5 \text{ tours.min}^{-1}$?

3. Tuba

1. Question préliminaire.

On considère une demi-sphère de rayon a plongée dans un liquide de masse volumique ρ (Cf. figure 34.15) et posée sur le fond du récipient. On note $P(0)$ la pression en O .

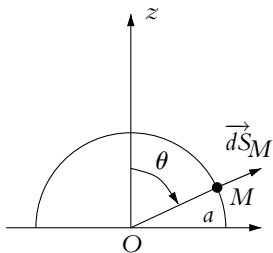


Figure 34.15 Question préliminaire

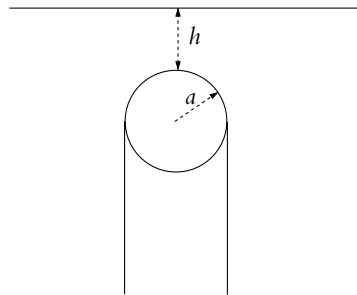


Figure 34.16 Question 2

- a) Exprimer la pression en M en fonction de $P(0)$, a et $\cos \theta$.
 - b) Comment est dirigée la résultante des forces de pression ? Exprimer la projection sur Oz de la force exercée sur une surface dS_M autour de M , en fonction de θ , $P(M)$ et dS_M .
 - c) Sachant que la surface élémentaire de largeur angulaire $d\theta$ faisant le tour de la demi-sphère est $dS = 2\pi a^2 \sin \theta d\theta$, calculer la force exercée sur la demi-sphère.
2. Un plongeur débutant utilise un tuba dont l'embout est fermé par une balle lorsqu'il plonge sous l'eau. La balle repose alors sur le tube (Cf. figure 34.16). On suppose que la moitié inférieure de la balle est dans le tube et la moitié supérieure en contact avec l'eau. Le rayon de la balle est noté $a = 1$ cm et, lors de la plongée, le sommet de la balle est à une profondeur $h = 10$ cm. On prendra $g = 10 \text{ m.s}^{-2}$ et la pression atmosphérique égale à $P_0 = 1$ bar.
- a) Déterminer la force exercée par l'eau sur la balle.
 - b) Pourquoi le plongeur doit-il souffler plus fort pour soulever la balle lorsqu'il plonge que lorsqu'il est dans l'air ?

4. Fermeture d'un bassin par une balle

Une sphère de bois de masse volumique ρ et de rayon R est complètement immergée dans un bassin d'eau (de masse volumique ρ_e) de profondeur H de telle manière qu'elle bouche un trou circulaire de rayon r .

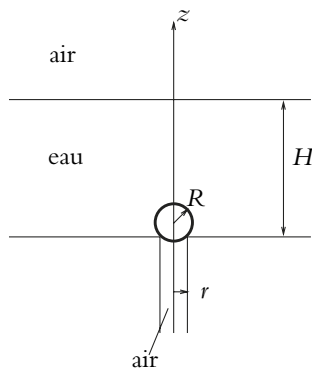


Figure 34.17

1. Calculer la force que la sphère exerce sur le fond du bassin.
2. Application numérique : $\rho = 850 \text{ kg.m}^{-3}$, $\rho_e = 1000 \text{ kg.m}^{-3}$, $H = 0,7 \text{ m}$, $R = 0,2 \text{ m}$, $r = 0,1 \text{ m}$ et $g = 9,8 \text{ m.s}^{-2}$.
3. Si on baisse le niveau de l'eau dans le bassin existe-t-il une hauteur d'eau pour laquelle la sphère monte à la surface avant qu'elle n'émerge ?

5. Étude d'un ballon sonde, d'après ENSTIM 2000

Si le premier ballon-sonde au dihydrogène est dû à Gustave Hermitte et Georges Besançon (1892), c'est incontestablement à l'ingéniosité et à la ténacité de l'atypique Léon Teisserenc de Bort (1855-1913) qu'on doit la mise au point des techniques d'investigation par ballon-sonde et la première cartographie atmosphérique.

On note (Oz) l'axe vertical ascendant, $z = 0$ au niveau du sol.

1. Etude de la troposphère.

La troposphère est la partie de l'atmosphère terrestre inférieure à 10 km. La troposphère fut dénommée ainsi en 1902 par Léon Teisserenc de Bort à partir de la racine "tropos", le changement. Il découvrit aussi la stratosphère. On la considère comme un gaz parfait de pression $p(z)$, de température $T(z)$ et de volume massique $v(z)$. Au sol, on a la pression p_0 et la température T_0 . Elle est en équilibre thermodynamique et mécanique et obéit à la loi polytropique empirique :

$$p^{-k}(z).T(z) = p_0^{-k}.T_0 \quad \text{avec} \quad k = 0,15 \quad (34.20)$$

- a) Comment peut-on qualifier la transformation correspondant au cas $k = 0$?
- b) Définir les mots "homogène" et "isotrope". Caractérisent-ils la troposphère ?
- c) Donner l'équation d'état d'un gaz parfait liant $p(z)$, $v(z)$, R , M_{air} et $T(z)$.
- d) Exprimer la loi de la statique des fluides avec g , $\frac{dp}{dz}$ et $v(z)$.
- e) On appelle gradient thermique la variation de la température par mètre :

$$\frac{dT(z)}{dz} = -\delta$$

Déduire δ en fonction de k , de la masse molaire de l'air M_{air} , de l'accélération de la pesanteur g et de la constante des gaz parfaits R . Calculer numériquement δ sachant que $M_{\text{air}} = 29 \text{ g.mol}^{-1}$, $g = 9,8 \text{ m.s}^{-2}$ et $R = 8,32 \text{ J.K}^{-1}.\text{mol}^{-1}$.

- f) Donner la loi de variation $T(z)$ en fonction de T_0 , δ et z .

g) On considère une quantité constante de n moles de gaz parfait à l'altitude z qui évolue dans la troposphère. On note $V(z)$ le volume qu'elle occupe à l'altitude z et V_0 son volume au sol. Déterminer la loi $\frac{V(z)}{V_0}$ en fonction de δ , z , T_0 et k .

2. Ascension d'un ballon sonde.

Le ballon sonde dégonflé et instrumenté a une masse totale $m_B = 1,2 \text{ kg}$. On gonfle au sol son enveloppe avec n_0 moles de dihydrogène. Son volume est alors V_0 . L'enveloppe reste fermée tant que son volume $V(z) < V_{\text{max}} = 10V_0$. Lorsque $V(z) = V_{\text{max}}$, l'enveloppe se déchire et le ballon retombe au sol.

a) Phase ascensionnelle à enveloppe hermétiquement fermée.

Sur ce ballon s'exerce une force de frottement $\vec{F}_f$. La force totale s'exerçant sur le ballon est $(F - m_{Bg})\vec{u}_z + \vec{F}_f$. En effectuant un bilan des forces, déterminer le terme F en fonction de n_0 , de g , de la masse molaire du dihydrogène $M_{H_2} = 2 \text{ g.mol}^{-1}$ et de celle de l'air M_{air} .

b) Calculer la valeur minimale $n_{\min}$ de n_0 pour que le ballon décolle.

c) On admet le modèle de troposphère précédent. Durant l'ascension, on peut considérer que la pression et la température sont quasiment identiques à l'intérieur et à l'extérieur du ballon. Calculer h , altitude maximale atteinte en prenant $T_0 = 293 \text{ K}$. Commenter le résultat.

6. Structure de l'atmosphère terrestre, d'après CCP PC 2000

L'atmosphère est essentiellement constituée d'un mélange gazeux comprenant du diazote pour 78% en volume, du dioxygène pour 21% en volume, de l'argon pour environ 1,0% en volume, du gaz carbonique pour 0,03% en volume et des traces infimes de néon, de krypton, d'hélium, d'ozone, de dihydrogène, de xénon et de divers rejets de la biosphère. Cette composition est à peu près constante jusqu'à une altitude de 85 km sauf pour certains gaz comme l'ozone surtout présent entre 30 et 40 km d'altitude.

L'atmosphère est stratifiée en température (la figure ci-dessous donne les variations de température en fonction de l'altitude) et donc aussi en pression.

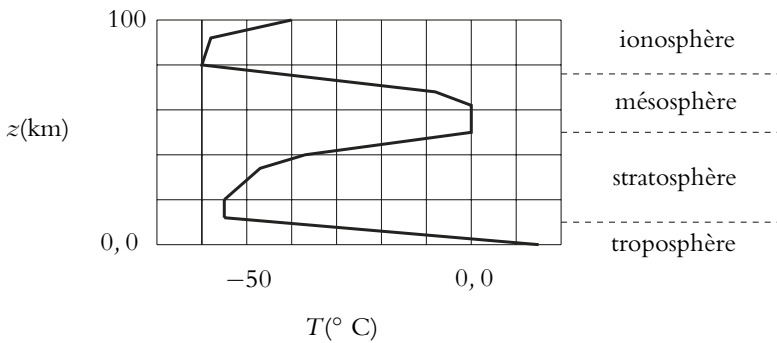


Figure 34.18

On considérera dans tout le problème que le mélange gazeux est sec au sens où on ne tiendra pas compte de la proportion variable de vapeur d'eau proportionnelle à la température et responsable des phénomènes de condensation en pluie, brouillard ou neige de l'air chaud qui s'élève et se refroidit.

On notera M la masse molaire de l'air, g le champ de pesanteur supposé uniforme, R la constante des gaz parfaits, k la constante de Boltzmann, n' le nombre de moles et n la densité volumique de molécules à l'altitude z .

On admettra que l'air est un gaz parfait.

1. Atmosphère isotherme :

On suppose dans un premier temps que la température est constante dans la troposphère.

- a) Établir que la pression ne dépend que de l'altitude z . On prendra l'axe vertical orienté vers le haut.
- b) Déterminer l'expression de la pression P en fonction de l'altitude en supposant que l'atmosphère est isotherme à la température T . On note P_0 la pression à l'altitude $z = 0$.
- c) Montrer que l'équation des gaz parfaits peut s'écrire sous la forme $P = nkT$.
- d) En déduire l'expression de n sous la forme $n_0 \exp\left(-\frac{E(z)}{kT}\right)$. On donnera les expressions de n_0 et $E(z)$.
- e) Donner la signification physique que $E(z)$ et kT .
- f) Calculer le rapport $\frac{P(z)}{P_0}$ pour $z = 10$ km en supposant une atmosphère isotherme à la température $T = 288$ K.
- g) Montrer que 70% de la masse totale de l'air est située en dessous de 10 km dans cette hypothèse.

2. Atmosphère adiabatique et polytropique :

On renonce à l'hypothèse d'une atmosphère isotherme pour passer à une atmosphère adiabatique.

- a) Les capacités thermiques molaires de l'air sont $C_V = \frac{5}{2}R$ et $C_P = \frac{7}{2}R$. Exprimer la valeur du coefficient γ .
- b) On rappelle que pour une transformation quasistatique adiabatique d'un gaz parfait, le produit PV^γ est constant. Montrer que le produit $T^\alpha P^\beta$ est constant et exprimer α et β en fonction de γ .
- c) En déduire la relation donnant $\frac{dT}{T}$ en fonction de $\frac{dP}{P}$ et γ .
- d) Établir l'expression du gradient de température adiabatique $\left(\frac{dT}{dz}\right)_{\text{adia}}$ en fonction de γ , M , g et R . Donner sa valeur pour l'air.
- e) À partir de la figure ci-dessus, donner la valeur de la température à 10 km ainsi que la valeur du gradient de température réel $\left(\frac{dT}{dz}\right)_{\text{réel}}$. La comparer à la valeur du gradient adiabatique.
- f) Les transformations réelles au sein de l'atmosphère ne sont ni isothermes ni adiabatiques mais entre les deux. On les dit polytropiques et elles vérifient PV^q constant avec $1 < q < \gamma$. Donner la valeur de q à partir des données expérimentales.
- g) En déduire la distribution correspondante de la température en fonction de l'altitude.
- h) Même question pour la pression en fonction de la température.
- i) Même question pour la pression en fonction de l'altitude.
- j) En déduire la pression et la température à 10 km.

35

Interprétation microscopique de T et P . Énergie interne. Cas du gaz parfait monoatomique

On va essayer d'interpréter physiquement les notions de température et de pression par une modélisation du gaz parfait monoatomique au niveau microscopique. Pour cela, il est nécessaire de préciser les hypothèses du modèle.

1. Les hypothèses

1.1 Propriétés du gaz

Pour modéliser le comportement d'un gaz, on suppose qu'il a les propriétés suivantes.

1. C'est un gaz parfait monoatomique, constitué d'un grand nombre N de particules **ponctuelles**.
2. Le gaz est à l'équilibre thermodynamique, la répartition est homogène et les variables d'état sont constantes dans le temps.
3. Les molécules sont en mouvement incessant et désordonné : on appelle ce mouvement *l'agitation thermique*.
4. Les seules interactions entre les particules sont les forces de contact au moment des chocs. Les chocs particule-particule et les chocs particule-paroi sont élastiques. On entend par choc élastique avec la paroi un choc qui n'est pas mou, c'est-à-dire que la particule n'est pas absorbée.

On pourrait penser que lors d'un choc, une particule peut céder une partie de son énergie cinétique qui servirait à exciter des niveaux d'énergie électronique de l'autre particule impliquée dans le choc. En effet, on sait que les niveaux d'énergie des atomes et des molécules sont quantifiés et que l'ordre de grandeur de l'énergie pour passer d'un niveau à l'autre est quelques électron-volts. Or pour qu'une particule ait une énergie cinétique de 1 eV, il faut que l'agitation soit suffisante et cela nécessite une température de 11 000 K non atteinte dans les systèmes thermodynamiques étudiés dans le cours.

► **Remarque :** Si les particules sont ponctuelles, il n'y a pas de chocs entre particules mais uniquement entre les particules et la paroi, mais on verra que la prise en compte des chocs entre particules est nécessaire pour expliquer la *thermalisation*.

Conséquences et remarques :

1. L'énergie mécanique du système se réduit à **l'énergie cinétique de translation**. Si les particules n'étaient pas considérées comme ponctuelles, il suffirait pour calculer leur énergie cinétique d'appliquer le théorème de Koenig qui donne pour un solide : $E_c = E_c^* + \frac{1}{2}mv_G^2$. L'énergie E_c^* correspond à une éventuelle rotation de la particule sur elle-même. Puisque on a supposé les particules ponctuelles, ce terme E_c^* est nul. D'autre part, dans l'hypothèse (4), on a supposé qu'il n'y avait pas d'interaction entre particules, donc l'énergie potentielle d'interaction est nulle. Cette énergie mécanique est constante globalement, mais pas individuellement (les particules peuvent échanger leur vitesse au moment du choc).
2. Toutes les particules n'ont pas la même vitesse. Cependant, si on considère un volume mésoscopique $d\tau$, il contient un très grand nombre de particules : on appelle *distribution des vitesses* l'ensemble $\{\vec{v}_i\}$ de toutes les vitesses des particules de $d\tau$.
 - L'équilibre thermodynamique impose que la **répartition statistique des vitesses soit la même à tout instant**.
 - En raison de l'homogénéité **elle est la même dans tout volume mésoscopique et dans l'ensemble du gaz**.
 - Les vitesses peuvent être dans n'importe quelle direction. On dit qu'il y a *isotropie* de l'espace, c'est-à-dire que toutes les directions sont équiprobables. Ainsi statistiquement, il y a autant de particules ayant une vitesse $\vec{v}$ que la vitesse opposée $-\vec{v}$ et :

$$\sum_{i=1}^N \vec{v}_i = 0 \quad (35.1)$$

1.2 Valeurs moyennes

Soit l'expérience des billes décrite au paragraphe 1.2.b) du chapitre 33. On a vu que lorsque le nombre de billes était très grand, ce qui correspond à un système thermodynamique, on ne pouvait avoir accès qu'à **la force moyenne** exercée par les billes sur la membrane, et non à la force exercée par une seule bille.

On peut généraliser cela à toutes les grandeurs d'un système thermodynamique, c'est-à-dire que seules les grandeurs moyennes sont accessibles. On rappelle la définition d'une grandeur moyenne : si on considère un ensemble de N particules et que, pour

chacune d'elle, la grandeur G a la valeur G_i , la valeur moyenne de G est définie par :

$$\overline{G} = \langle G \rangle = \frac{1}{N} \sum_{i=1}^N G_i \quad (35.2)$$

On peut appliquer cette définition pour calculer la vitesse moyenne $\langle \vec{v} \rangle$. En raison de l'isotropie de l'espace, on obtient avec la relation (35.1) :

$$\langle \vec{v} \rangle = \frac{1}{N} \sum_{i=1}^N \vec{v}_i = \vec{0}$$

Par contre, les valeurs moyennes de la norme de la vitesse v et de la vitesse au carré ne sont pas nulles, puisqu'elles résultent de la somme de termes positifs.

On définit une grandeur importante qui intervient dans la suite, la *vitesse quadratique moyenne* u qui est la racine carrée de la valeur moyenne du carré des vitesses :

$$u = \sqrt{\langle v^2 \rangle} = \sqrt{\frac{1}{N} \sum_{i=1}^N v_i^2} \quad (35.3)$$



Attention, $\langle v^2 \rangle \neq \langle v \rangle^2$.

On s'intéresse maintenant aux valeurs moyennes des composantes de $\vec{v}$: v_x, v_y, v_z . Par isotropie, la distribution des vitesses est la même dans un sens ou l'autre pour chaque direction donc pour toute particule ayant une composante v_x , il y a une particule ayant la composante opposée soit :

$$\langle v_x \rangle = 0$$

Par le même raisonnement :

$$\langle v_y \rangle = 0 \quad \text{et} \quad \langle v_z \rangle = 0$$

Mais comme pour v^2 , les carrés des valeurs moyennes des composantes ne sont pas nulles. On peut les exprimer en fonction de la vitesse quadratique moyenne u . En effet, les trois directions de l'espace sont équivalentes donc les valeurs moyennes suivant ces trois directions sont égales et on peut écrire :

$$\langle v_x^2 \rangle = \langle v_y^2 \rangle = \langle v_z^2 \rangle$$

Sachant que :

$$u^2 = \langle v^2 \rangle = \langle v_x^2 \rangle + \langle v_y^2 \rangle + \langle v_z^2 \rangle$$

on obtient :

$$\langle v_x^2 \rangle = \langle v_y^2 \rangle = \langle v_z^2 \rangle = \frac{1}{3} \langle v^2 \rangle = \frac{1}{3} u^2 \quad (35.4)$$

2. Calcul de la pression

2.1 Présentation du modèle

Le gaz considéré occupe un volume V , à la température T et contient n particules par unité de volume (Cf. figure 35.1) de masse m .

Pour calculer la pression exercée sur la paroi, il est nécessaire de calculer la force exercée par les particules du gaz sur un élément de surface dS de cette paroi soit $d\vec{F}_{\text{particules} \rightarrow dS}$. Si on applique la relation fondamentale de la dynamique à cet élément de paroi, cette force est égale à la variation de la quantité de mouvement de dS (Cf. figure 35.1), due aux chocs des particules, sur l'intervalle de temps dt , soit :

$$d\vec{p}_{dS} = d\vec{F}_{\text{particules} \rightarrow dS} dt$$

Or, d'après le principe des actions réciproques, la force exercées par les particules sur la paroi est l'opposé de la force exercée par la paroi sur les particules, donc la variation de la quantité de mouvement de l'élément de paroi dS due au choc des particules est l'opposé de la variation de quantité de mouvement des particules venant frapper dS , soit :

$$d\vec{p}_{dS} = -d\vec{p}_{\text{particules}, dS} \quad (35.5)$$

Pour déterminer $d\vec{F}_{\text{particules} \rightarrow dS}$, il s'agit de calculer cette variation de quantité de mouvement $d\vec{p}_{\text{particules}, dS}$.

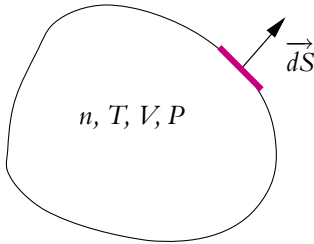


Figure 35.1 Éléments de surface.

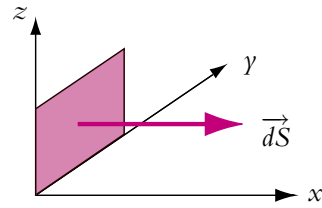


Figure 35.2 Choix du repère.

On choisit les axes Ox , Oy et Oz (vecteurs directeurs $\vec{u}_x$, $\vec{u}_y$, $\vec{u}_z$) tels que $d\vec{S}$ soit suivant Ox , c'est-à-dire : $d\vec{S} = dS\vec{u}_x$. (Cf. figure 35.2).

Pour le calcul de la pression exercée par le gaz sur son enveloppe, on effectue les hypothèses simplificatrices suivantes :

1. les particules ont toutes la même vitesse en module,
2. la vitesse commune est égale à la vitesse quadratique moyenne u ,
3. chaque particule ne se déplace que suivant une seule des trois directions Ox , Oy ou Oz , et suivant un seul sens,

4. en raison de l'isotropie de l'espace, les six possibilités de déplacement sont équiprobables et il y a **un sixième des particules** qui se déplacent suivant chaque sens.

► **Remarque :** Une modélisation plus compliquée, qui n'est pas au programme, permet de justifier ces trois hypothèses.

2.2 Modèle simplifié avec choc frontal

On calcule tout d'abord la variation de quantité de mouvement de dS due au choc d'une particule sur la paroi. Seules les particules ayant une vitesse initiale $\vec{v}$ suivant $\vec{u}_x$ peuvent rencontrer la paroi. Étant donné qu'elles rencontrent la paroi sous incidence normale, elles repartent après le choc avec une vitesse finale $\vec{v}'$ suivant $-\vec{u}_x$. La variation de la quantité de mouvement d'une de ces particules est donc :

$$\vec{dp}_{\text{particule}} = m(\vec{v}' - \vec{v}) = m(-u\vec{u}_x - u\vec{u}_x) = -2mu\vec{u}_x \quad (35.6)$$

Or comme cela a été signalé, en raison du choc, la variation de la quantité de mouvement de la paroi due au choc de cette particule est l'opposé de la variation de quantité de mouvement de la particule, soit :

$$\vec{dp}_{\text{paroi}} = -\vec{dp}_{\text{particule}} = 2mu\vec{u}_x \quad (35.7)$$

Il faut maintenant déterminer combien de particules viennent heurter la paroi pendant un intervalle de temps dt . Dans cet intervalle, les particules auront parcouru la distance $dL = u dt$, si elles ne rencontrent aucun obstacle. Ainsi, les particules venant heurter la paroi pendant dt doivent se situer à une distance inférieure ou égale à dL . Ce sont donc les particules occupant le volume parallélépipédique $d\tau$ de longueur dL , de section dS situé devant la surface dS (Cf. figure 35.3).

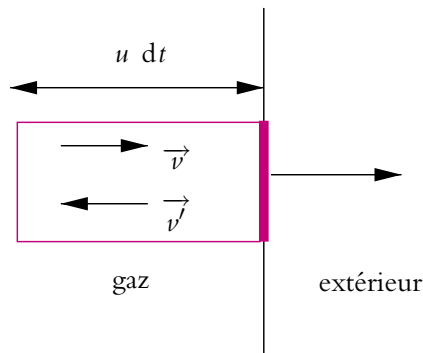


Figure 35.3 Parallélépipède contenant les particules qui viennent frapper dS pendant dt .

Ce volume contient $n d\tau$ particules dont seulement $1/6$ se dirigent vers la paroi. Le nombre de particules entrant en collision avec $d\vec{S}$ pendant dt est donc :

$$dN = \frac{1}{6} n d\tau = \frac{1}{6} n dS u dt \quad (35.8)$$

Chaque particule de l'équation (35.8) entraînant une variation de la quantité de mouvement de dS donnée par (35.7), la variation totale de quantité de mouvement de dS pendant dt est :

$$d\vec{p} = dN \times 2mu\vec{u}_x = \frac{1}{3} nm u^2 dS dt \vec{u}_x \quad (35.9)$$

Sachant que la force $d\vec{F}$ exercée sur dS vérifie : $d\vec{p} = d\vec{F} dt$, on en déduit :

$$d\vec{F} = \frac{1}{3} nm u^2 dS \vec{u}_x$$

et puisque $d\vec{F} = P d\vec{S}$, alors la pression est donnée par :

$$P = \frac{1}{3} nm u^2 \quad (35.10)$$

La pression est donc proportionnelle à la densité particulaire, à la masse des particules et à la vitesse quadratique moyenne au carré.

La pression calculée ainsi porte le nom de *pression cinétique*.

3. Température cinétique

Si le modèle précédent est bon, on doit avoir égalité entre la pression cinétique de l'expression (35.10) et la pression figurant dans l'équation du gaz parfait. Ainsi, sachant que $n = \frac{N}{V}$:

$$P = \frac{1}{3} nm u^2 = nk_B T$$

On en déduit une expression de la température appelée *température cinétique* :

$$T = \frac{mu^2}{3k_B} \quad (35.11)$$

On obtient aussi l'expression de la vitesse quadratique moyenne en fonction de la température :

$$u = \sqrt{\frac{3k_B T}{m}} = \sqrt{\frac{3RT}{M}} \quad (35.12)$$

où M est la masse molaire du gaz et m la masse d'une molécule.

Les deux expressions précédentes sont très importantes. Elles montrent que la température est directement liée au mouvement des particules. Plus la température est grande, plus la vitesse quadratique moyenne est grande, donc plus les particules sont « agitées » : c'est ce qu'on appelle *l'agitation thermique*. Ainsi, les déplacements (et les chocs) des particules à l'échelle microscopique, qui sont un phénomène purement mécanique, sont ressentis à l'échelle macroscopique par l'intermédiaire d'une grandeur (la température) qui n'existe pas en mécanique. On peut calculer u dans le cas de deux gaz pour se rendre compte de l'ordre de grandeur important de cette vitesse à une température de 300 K :

- pour H_2 , $u = \sqrt{\frac{3 \times 1.38 \cdot 10^{-23} \times 300}{2 \times 1.67 \cdot 10^{-27}}} = 1,9 \cdot 10^3 \text{ m} \cdot \text{s}^{-1}$,
- pour O_2 , $u = 465 \text{ m} \cdot \text{s}^{-1}$.

On peut dire que « ça décoiffe » ! Les vitesses des particules peuvent donc être très importantes. On rappelle que u est une moyenne, donc certaines particules sont plus rapides et d'autres plus lentes. Ces vitesses étant importantes, on peut penser que les particules parcourent des distances importantes ; c'est oublier que les chocs sont nombreux (en réalité, les particules ne sont pas ponctuelles, donc il existe des chocs entre particules). Ce sont d'ailleurs ces chocs qui permettent d'assurer la répartition des vitesses et l'homogénéisation de la température : on dit alors que le système est *thermalisé*.

4. Énergie interne

Comme son nom l'indique, c'est l'énergie « contenue » dans le système, on la note U . De manière générale, elle est la somme de plusieurs termes :

- l'énergie cinétique microscopique de translation des particules,
- l'énergie cinétique de rotation (sur elles-mêmes) des particules non ponctuelles,
- l'énergie de vibration des molécules (modélisation des liaisons par des ressorts vue en mécanique),
- l'énergie potentielle d'interaction entre particules (la force dont dérive cette énergie potentielle est la force de Van der Waals).

4.1 Cas du gaz parfait monoatomique

Pour le gaz parfait monoatomique, comme les particules sont supposées ponctuelles et qu'il n'y a pas de forces d'interaction, seule l'énergie cinétique microscopique de translation est à prendre en compte. Ainsi, pour les N particules, on peut calculer l'énergie interne U :

$$U = \sum_{i=1}^N \frac{1}{2} m v_i^2 = \frac{1}{2} m \sum_{i=1}^N v_i^2$$

or sachant que $u = \frac{1}{N} \sum v_i^2$, on obtient :

$$U = \frac{1}{2} m N u^2 \Rightarrow U = \frac{3}{2} N k T = \frac{3}{2} n_m R T \quad (35.13)$$

où n_m est le nombre de moles du gaz. L'énergie interne du gaz parfait monoatomique est proportionnelle à la température.

On définit la *capacité thermique à volume constant* C_v comme la quantité d'énergie nécessaire pour élever la température d'un degré à volume constant, soit pour le gaz parfait monoatomique :

$$C_v = \frac{3}{2} N k = \frac{3}{2} n_m R \quad (35.14)$$

L'unité de C_v est le J.K^{-1} et l'énergie interne du gaz parfait monoatomique est :

$$U = C_v T \quad (35.15)$$

L'origine du nom sera vue au paragraphe (4.4). On peut définir aussi la capacité thermique molaire à volume constant $C_{vm} = \frac{C_v}{n_m}$ en $\text{J.K}^{-1}.\text{mol}^{-1}$ et la capacité thermique massique à volume constant $c_v = \frac{C_{vm}}{M} = \frac{C_v}{m}$, (M étant la masse molaire et m la masse du gaz étudié) en $\text{J.K}^{-1}.\text{kg}^{-1}$. Pour un gaz parfait monoatomique, $C_{vm} = 12,5 \text{ J.K}^{-1}.\text{mol}^{-1}$, qui est donc l'énergie nécessaire pour élever la température d'une mole de gaz d'un degré.

Attention, les notations pour les capacités thermiques ne sont pas standardisées. Le plus simple est de regarder l'unité pour savoir de quelle capacité il s'agit.

4.2 Extension aux gaz parfaits polyatomiques

Dans le cas du gaz parfait polyatomique, la molécule ne peut plus être considérée ponctuelle. Il faut donc ajouter à l'énergie cinétique de translation celle de rotation de la molécule sur elle-même. Cette énergie est aussi proportionnelle à la température T et l'expression (35.15) est toujours valable, cependant l'expression de C_v est différente. Ainsi pour un gaz parfait diatomique, on trouve :

$$U = C_v T \quad \text{avec} \quad C_v = \frac{5}{2} n_m R = \frac{5}{2} N k_B \quad (35.16)$$

Pour le gaz parfait diatomique, la capacité thermique molaire à volume constant est : $C_{vm} = 20,8 \text{ J.K}^{-1}.\text{mol}^{-1}$. Il faut donc plus d'énergie pour augmenter la température d'un gaz parfait diatomique d'un kelvin que pour un gaz parfait monoatomique.

► **Remarque :** Le terme en facteur de U est $5/2$ pour le gaz diatomique au lieu de $3/2$ pour le gaz monoatomique car une molécule diatomique possède cinq degrés de liberté

(3 de translation et 2 de rotation) contrairement à la molécule monoatomique qui n'en possède que 3. On reviendra sur ce point dans le paragraphe suivant.

On ne donne pas d'expression générale de la capacité thermique à volume constant C_v , mais pour un gaz parfait, on peut la considérer constante dans un très large intervalle de température et l'énergie interne U proportionnelle à la température T , soit :

$$U = C_v T \quad \text{avec } C_v \text{ constant} \quad (35.17)$$

Un gaz parfait obéit à la *première loi de Joule*, c'est-à-dire que son énergie interne ne dépend que de la température.

Pour une élévation de la température de dT , la variation élémentaire de l'énergie interne est :

$$dU = C_v dT \quad (35.18)$$

4.3 Interprétation de l'expression de l'énergie interne

Sans établir de théorie générale, on va essayer d'expliquer pourquoi un gaz polyatomique nécessite plus d'énergie pour être chauffé qu'un gaz monoatomique. Il faut pour cela introduire la notion de *degrés de liberté*. Pour une particule ponctuelle, il existe trois degrés de liberté : les trois directions de l'espace suivant lesquelles elle peut se déplacer. Les trois directions sont équiprobables et donc elles nécessitent chacune autant d'énergie soit $1/3$ de l'énergie interne ou $\frac{1}{2}NkT$. Statistiquement pour chaque

particule, un degré de liberté va donc nécessiter $\frac{1}{3} \frac{U}{N} = \frac{1}{2}kT$.

Une molécule diatomique possède plus de degrés de liberté car comme cela a été vu en mécanique avec les théorèmes de Koenig, le mouvement de la molécule se décompose en un mouvement de translation du centre de gravité (trois degrés de liberté) composé à un mouvement de rotation autour de son centre de gravité. Si on imagine la molécule rigide, comme une règle, il y a deux possibilités de rotation autour des deux axes de rotation perpendiculaires à la direction de la liaison. Cela rajoute deux degrés de liberté aux trois degrés de translation soit **cinq degrés** de liberté en tout, eux aussi **équiprobables**. Si l'on compte comme précédemment, statistiquement $\frac{1}{2}kT$ par molécule et par degré de liberté cela fait en tout une énergie $\frac{5}{2}kT$ par molécule. On retrouve l'expression (35.16) pour l'énergie interne du gaz, soit $U = \frac{5}{2}NkT$.

Ainsi l'énergie s'homogénéise entre les différents degrés de liberté du système. On parle d'*équpartition de l'énergie*. Chaque degré de liberté « consomme » $\frac{1}{2}NkT$. Plus le nombre d'atomes augmente dans la molécule, plus le nombre de degrés de liberté augmente et plus il faut d'énergie pour pouvoir tous les exciter.

On revient à l'étude de la capacité thermique du gaz parfait diatomique. Si on augmente fortement la température, la molécule diatomique va se mettre à vibrer, la liaison se comportant comme un ressort. La figure 35.4 récapitule l'évolution de la capacité thermique d'un gaz parfait diatomique. La température nécessaire pour que la molécule tourne sur elle-même est de quelques kelvins (2 K pour O_2) à quelques dizaines de kelvins (85 K pour H_2). Par contre, pour qu'elle se mette à vibrer, cela nécessite une température beaucoup plus importante : de l'ordre de 2200 K pour O_2 et 6000 K pour H_2 . On remarque que l'énergie correspondant à la vibration est kT et non $\frac{1}{2}kT$. Il se rajoute donc deux degrés de liberté. En effet, contrairement à la translation et à la rotation pour lesquelles il n'y a que de l'énergie cinétique, la vibration met en œuvre une énergie cinétique et une énergie potentielle du type ressort.

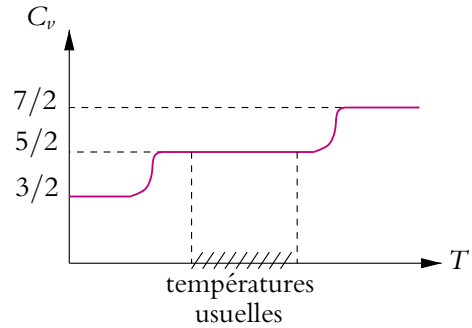


Figure 35.4 Évolution de la capacité thermique d'un gaz diatomique avec T .

Aux températures usuelles, les molécules diatomiques se comportent donc comme des molécules rigides.

À très forte température, ce sont les transitions électroniques qui peuvent avoir lieu. La capacité thermique est donc plus importante. Mais, comme cela a déjà été signalé, il faut une température de 11000 kelvins par électron-volt pour une transition électronique. À ces températures, la plupart des molécules sont détruites.

4.4 Cas du gaz de Van der Waals

Dans le cas plus réaliste du gaz de Van der Waals, on a vu qu'on tenait compte de l'interaction entre molécules. Il faut donc ajouter à l'expression de l'énergie interne obtenue pour le gaz parfait une énergie potentielle d'interaction. On admet que l'énergie interne d'un gaz de Van der Waals est :

$$U = U_{GP}(T) - \frac{n_m^2 a}{V} \quad (35.19)$$

où $U_{GP}(T)$ est l'énergie interne du gaz parfait correspondant (aux faibles pressions, tout gaz se comporte comme un gaz parfait) et a la constante définie dans l'équation du gaz de Van der Waals (33.9). Le terme ajouté à l'énergie interne du gaz parfait est négatif. En effet, ce terme correspond à l'énergie potentielle d'attraction entre les molécules. Si on se réfère au cours de mécanique, l'énergie potentielle d'une force attractive qui assure une cohésion du système est négative (on choisit l'énergie poten-

tielle nulle quand les particules sont infiniment éloignées les unes des autres, donc sans interaction).

On remarque que l'énergie interne d'un gaz non parfait dépend de la température et du volume. C'est une propriété générale qu'on admet. L'énergie interne d'un gaz est une fonction de T et V : $U(T, V)$. Dans le cas du gaz parfait, l'énergie interne ne dépend que de T .

On définit la capacité thermique à volume constant C_v comme la dérivée de U par rapport à T à volume constant (d'où le nom) :

$$C_v = \left(\frac{\partial U}{\partial T} \right)_V \quad (35.20)$$

La définition précédente est applicable aussi pour un gaz parfait même si l'expression de U ne dépend pas du volume. Dans ce cas on a :

$$C_v = \frac{dU}{dT}$$

4.5 Extension aux solides et liquides peu compressibles

Dans ce cas aussi, l'énergie interne ne dépend que de la température et du volume. Cependant les solides et liquides peu compressibles sont **très peu sensibles aux variations de volume**. On peut donc négliger l'effet d'une variation de volume et admettre que l'énergie interne ne dépend pratiquement que de la température. Ceci est d'ailleurs bien vérifié expérimentalement. Dans ce cas, pour une variation élémentaire dT de la température, la variation élémentaire de l'énergie interne dU :

$$dU = C_v dT \quad (35.21)$$

C_v dépend de la température mais peut être considérée comme constante dans un large domaine de température. Attention, même si l'expression générale de dU est la même que pour un gaz parfait, la capacité thermique à volume constant n'a pas la même valeur que pour un gaz parfait, elle est beaucoup plus grande (par exemple, pour l'eau, $c_v = 4,2 \text{ J.K}^{-1}.\text{g}^{-1}$, soit $C_{vm} = 75,2 \text{ J.K}^{-1}.\text{mol}^{-1}$). Si on peut considérer C_v constante dans le domaine de températures étudié, la variation d'énergie interne d'un solide ou d'un liquide, entre deux états de température T_1 et T_2 , est :

$$\Delta U = U_2 - U_1 = C_v(T_2 - T_1) \quad (35.22)$$

A. Applications directes du cours

1. Existence d'une atmosphère à la surface des planètes

1. Calculer numériquement la vitesse de libération (Cf. cours de mécanique sur les mouvements newtoniens) et la vitesse quadratique moyenne pour le dihydrogène et pour le diazote à la surface des quatre planètes telluriques (Mercure, Vénus, Terre, Mars), pour une température de 300 K. Conclure quant à la présence possible d'une atmosphère à cette température. On donne :

Planète	Diamètre équatorial en km	Rapport de la masse à celle de la Terre
Mercure	4878	0,055
Vénus	12 104	0,815
Terre	12 756	1
Mars	6794	0,107

Masse de la Terre : $M_T = 5,976 \cdot 10^{24}$ kg ; $G = 6,67 \cdot 10^{-11}$ m³kg⁻¹s⁻² ; masse des molécules : $m_{H_2} = 3,3 \cdot 10^{-27}$ kg et $m_{N_2} = 4,65 \cdot 10^{-26}$ kg ; constante Boltzmann : $k_B = 1,38 \cdot 10^{-23}$ J.K⁻¹.

2. Quel devrait être l'ordre de grandeur de la température pour que les molécules de diazote échappent à l'attraction terrestre ?

2. Effusion d'un gaz

On considère deux compartiments de volume V_1 et V_2 , l'ensemble est maintenu à la température T . Entre les deux compartiments, un petit trou de section s a été pratiqué. Initialement on a N_a particules d'un gaz parfait dans V_1 . On note N_1 et N_2 les nombres de particules dans les volumes V_1 et V_2 et on adopte le modèle suivant pour lequel les particules ont toutes le même module de vitesse v et que leur vitesse est suivant l'une des directions $\vec{u}_x$, $\vec{u}_y$, $\vec{u}_z$, $-\vec{u}_x$, $-\vec{u}_y$, $-\vec{u}_z$.

1. Quel est le nombre $dN_{1 \rightarrow 2}$ passant de V_1 à V_2 entre t et $t + dt$?
2. En déduire les équations différentielles vérifiées par N_1 et N_2 en fonction de N_1 , N_2 , s , v et $V = V_1 = V_2$.
3. Établir les expressions de N_1 et de N_2 en fonction du temps.
4. Définir un temps caractéristique τ .
5. Comment varie-t-il en fonction de la masse du gaz si on admet que $v = \sqrt{\frac{3RT}{M}}$?
6. Quelle peut être l'application pratique de ce phénomène d'effusion gazeuse ?

B. Exercices et problèmes

1. Force exercée sur un piston, d'après Agregation interne 2002.

On étudie l'enceinte de la figure ci-contre constituée d'un cylindre de section S fermé par un piston coulissant sans frottement. L'axe du cylindre est selon l'axe Ox , de vecteur unitaire $\vec{u}_x$.

On note L la distance entre le fond du cylindre et le piston. L'enceinte contient des particules assimilées à des points matériels de même masse μ . Par hypothèse, ces particules se déplacent uniquement avec une vitesse $\vec{v}$ suivant $\vec{u}_x$.

Elles subissent des chocs alternativement sur le fond du cylindre et sur le piston.

Le piston est fixe.

Les chocs des particules sur la paroi et le piston sont élastiques, c'est-à-dire que l'énergie cinétique totale est conservée au cours d'un choc.

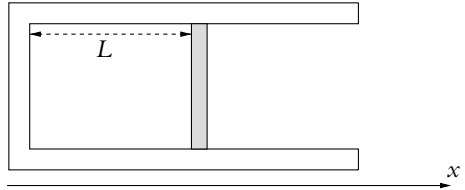


Figure 35.5

1. Montrer que la norme de la vitesse d'une particule se conserve au cours d'un choc.
2. Exprimer, en fonction de μ , L et v , la variation algébrique Δp_x de la quantité de mouvement de la particule au cours d'un choc sur le piston, et l'intervalle de temps Δt entre deux chocs successifs sur le piston.
3. Représenter schématiquement les variations de la force instantanée exercée par la particule sur le piston.
4. On définit l'intervalle de temps $\tau = n\Delta t$. Quelle est la variation de la quantité de mouvement $\Delta p'_x(\tau)$ du piston pendant l'intervalle τ , en fonction de n , μ et v ? En déduire la force moyenne exercée F_m sur le piston :

$$F_m = \lim_{n \rightarrow \infty} \frac{\Delta p'_x(\tau)}{\tau}$$

5. L'enceinte contient maintenant N particules identiques, de même masse μ mais dont les vitesses ont des normes v_i ($1 \leq i \leq N$) *a priori* différentes. Elles n'interagissent pas les unes avec les autres. On note :

$$\overline{E}_c = \frac{1}{N} \sum_{i=1}^N \frac{1}{2} \mu v_i^2$$

l'énergie cinétique moyenne des particules. Calculer, en fonction de L , N et $\overline{E}_c$, la force moyenne exercée par les particules sur le piston.

6. On définit la température T de l'ensemble des particules par la relation :

$$\overline{E}_c = \frac{1}{2} k_B T$$

où k_B est la constante de Boltzmann. Montrer qu'on retrouve pour le gaz de particules l'équation d'état du gaz parfait.

Premier principe

36

1. Premier principe – Énergie interne

1.1 Énergie interne – Énergie totale

On considère un système thermodynamique **macroscopiquement au repos** (c'est-à-dire qu'il n'y a pas de mouvement macroscopique interne, ni de mouvement d'ensemble). D'après le chapitre précédent :

L'énergie interne U est la somme de l'énergie cinétique microscopique totale E_{CM} et de l'énergie potentielle totale d'interaction entre les particules du système E_{PM} : $U = E_{PM} + E_{CM}$.

Exemple

Dans le cas du gaz parfait monoatomique, $E_{CM} = \frac{3}{2}nRT$ et $E_{PM} = 0$ soit $U = \frac{3}{2}nRT$. En fait, comme toute énergie, U est définie à une constante près, on peut donc écrire $U = \frac{3}{2}nRT + U_0$ où U_0 contient toute les formes d'énergie ne participant pas aux transformations du système (énergie atomique, nucléaire).

L'énergie interne représente-t-elle l'énergie totale E du système ? C'est le cas dans de nombreuses situations mais les trois exemples suivants montrent qu'il n'y a pas toujours identité entre énergie interne et énergie totale.

- Le système est globalement en mouvement à la vitesse V . Dans ce cas, il suffit d'appliquer le théorème de Koenig ($E_c = E_c^* + \frac{1}{2}MV^2$), E_c^* l'énergie cinétique dans le référentiel barycentrique étant comptabilisée dans l'énergie interne et M étant la masse totale du système, l'énergie totale s'écrit : $E = U + \frac{1}{2}MV^2$
- Lorsque certaines parties de masse m_J du système sont en mouvement à la vitesse V_J (par exemple, un piston), l'énergie totale s'écrit : $E = U + \frac{1}{2} \sum_J m_J V_J^2$

- Lorsque les forces extérieures macroscopiques s'exerçant sur le système dérivent d'une énergie potentielle E_p l'énergie totale s'écrit : $E = U + E_p$.

Dans le cas le plus général :

$$E = U + E_c + E_p$$

où E_p est l'énergie potentielle associée aux forces extérieures et E_c l'énergie cinétique macroscopique du système, c'est-à-dire $\frac{1}{2}MV^2$ si le système est animé d'un mouvement global à la vitesse V ou $\frac{1}{2}\sum_j m_j V_j^2$ si certaines parties de masse m_j du système sont en mouvement à la vitesse V_j .

Propriétés de l'énergie interne :

- L'état du système est défini par les variables d'état ($P, V, T, \dots$). L'énergie interne (ou totale) est une fonction de ces variables d'état : c'est une *fonction d'état*.
- L'énergie interne est une grandeur extensive.
Pourtant elle contient l'énergie potentielle d'interaction entre particules qui, comme on l'a vu en mécanique, n'est pas une grandeur additive : l'énergie potentielle d'interaction de N particules n'est pas égale à N fois l'énergie d'interaction entre deux particules. L'extensivité de l'énergie potentielle microscopique est due au fait que les forces d'interaction (forces de Van Der Waals) sont de très courte portée et que deux volumes mésoscopiques voisins n'ont pas d'interaction entre eux.



À un état macroscopique correspond une valeur de l'énergie interne et une seule (une fois la constante U_0 déterminée), mais à une valeur de l'énergie interne peuvent correspondre plusieurs états d'équilibre (par exemple, pour le gaz parfait, on peut avoir des états à la même température mais de pression et volume différents).

1.2 Énoncé du premier principe de la thermodynamique

En thermodynamique, on peut rencontrer deux types de système :

- Les systèmes *fermés* qui n'échangent pas de matière avec l'extérieur ;
- Les systèmes *ouverts* qui échangent de la matière avec l'extérieur.

En général, les systèmes étudiés seront des systèmes fermés.

a) Énoncé du premier principe

À tout système thermodynamique fermé (S) est associée une fonction d'état E appelée *énergie*. Au cours d'une transformation quelconque, la variation de E est égale à l'énergie reçue (positive si elle effectivement reçue, négative si elle est effectivement cédée) par le système de la part du milieu extérieur.

Le premier principe est donc un principe de conservation.

Le premier principe s'écrit :

$$E_{\text{finale}} - E_{\text{initiale}} = \mathcal{E}_{\text{recue}} \text{ avec } E = U + E_c + E_p$$

où l'énergie reçue $\mathcal{E}_{\text{recue}}$ est positive si elle est effectivement reçue par le système, et négative si elle est effectivement fournie par le système.

On note généralement :

$$\Delta E = E_{\text{finale}} - E_{\text{initiale}}$$

Dans de nombreux cas (système globalement au repos ou énergie cinétique macroscopique négligeable et sans variation d'énergie potentielle)

$$\Delta E = \Delta(U + E_c + E_p) = \Delta U$$

Si on plonge une bouteille d'eau à 300 K dans une baignoire d'eau chaude à 350 K, la température de l'eau de la bouteille augmente. Il n'y a pourtant aucun travail échangé puisque les systèmes sont immobiles. Cela prouve qu'il existe une autre forme d'échange d'énergie que le travail mécanique : on appelle cet échange d'énergie le *transfert thermique* (noté Q).

L'énergie reçue se décompose en deux formes d'énergie qu'on appelle *travail* W (c'est le travail défini en mécanique) et *transfert thermique*¹ Q (terme supplémentaire par rapport à ce qui a été vu en mécanique).

L'expression du premier principe est donc :

$$\begin{aligned} \Delta E &= W + Q \text{ soit } \Delta(U + E_c + E_p) = W + Q \\ \text{ou, si } E &= U, \Delta U = W + Q \end{aligned} \quad (36.1)$$

Le premier principe s'applique entre deux états d'équilibre.

► Remarques

- On en déduit que, pour un système isolé, l'énergie totale E est constante. Attention, il y a une différence de définition entre un système isolé en thermodynamique qui n'échange pas d'énergie avec l'extérieur et un système isolé en mécanique qui ne subit pas d'action extérieure.
- On s'aperçoit que par rapport à la mécanique, il y a un terme supplémentaire dans le bilan d'énergie : le transfert thermique.

¹ On rencontre aussi le terme quantité de chaleur.

1.3 Convention d'orientation du système

Pour un système (S) donné, comment orienter le travail et le transfert thermique reçu par le système? La première chose à faire est de **bien définir le système** (voir paragraphe 1.4). Ensuite, comme en électricité, on choisit de compter positivement l'énergie effectivement reçue par le système.

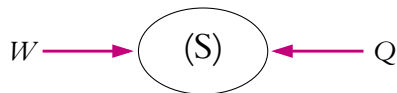


Figure 36.1 Si W (ou Q) est effectivement reçu par le système, W (ou Q) est positif. Si W (ou Q) est effectivement cédé par le système, W (ou Q) est négatif.

1.4 Délimitation du système

Il faut que les échanges de travail et de transfert thermique du système considéré soient bien définis et donc que les limites du système soient parfaitement connues.

En thermodynamique comme dans tout problème physique, il faut avant tout définir le système étudié.

On prend l'exemple d'un corps A en translation à la vitesse $\vec{V}$, glissant avec frottement $\vec{T}$ sur un support fixe B , sous l'action d'une force $\vec{F}$ extérieure à A et B qui fournit un travail W entre les instants t_1 et t_2 (Cf. figure 36.2)

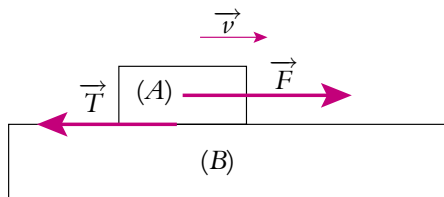


Figure 36.2 Délimitation du système.

En mécanique, on raisonne sur A puisqu'on connaît les forces qui s'exerce sur ce système. En thermodynamique, on est obligé de raisonner sur le système $A + B$ car on ne sait pas comment se fait le transfert thermique dû au frottement (de A vers B ou de B vers A ?).

On applique le premier principe au système $A + B$ entre t_1 et t_2 :

$$\left(\frac{1}{2}mv_2^2 + U_A(t_2) + U_B(t_2) \right) - \left(\frac{1}{2}mv_1^2 + U_A(t_1) + U_B(t_1) \right) = W + Q \quad (36.2)$$

Le travail de la force de frottement n'intervient pas dans le premier principe car c'est une force intérieure au système (voir énoncé paragraphe 1.2).

On applique le théorème de l'énergie cinétique à A en appelant W_T le travail de la force de frottement

$$\frac{1}{2}mv_2^2 - \frac{1}{2}mv_1^2 = W + W_T \quad (36.3)$$

En combinant les équations (36.2) et (36.3) on obtient :

$$(U_A(t_2) + U_B(t_2)) - (U_A(t_1) + U_B(t_1)) = Q - W_T \quad (36.4)$$

Si le système est isolé thermiquement, $Q = 0$. Or le travail des forces de frottements est résistant donc $W_T < 0$. Dans ce cas, l'énergie interne des deux corps augmente, et puisque ce sont deux solides, leur température augmente aussi (en effet, pour un solide, $U = C T + \text{cste}$).

Si le système est maintenu à température constante, l'énergie interne U est constante aussi et $Q = W_T < 0$: le système cède du transfert thermique à l'extérieur.

Il est important de constater sur ces deux exemples qu'une variation de température n'est pas forcément associée à un transfert thermique.

2. Transformations d'un système thermodynamique

2.1 Cas général

Soit un système (S) évoluant d'un état initial (i) à un état final (f) ; la variation d'énergie interne entre ces deux états est $\Delta U = U_f - U_i = W + Q$.

La variation d'énergie interne ne dépend que des états extrêmes et non du chemin suivi par la transformation (on retrouve cette propriété pour l'énergie potentielle en mécanique) : cette propriété est caractéristique d'une fonction d'état.

La somme $W + Q$ ne dépend pas du chemin suivi mais *a priori* W et Q individuellement dépendent des états extrêmes, mais aussi des transformations entre l'état initial et l'état final (voir exemples ultérieurs).

On distingue différents types de transformations parmi elles on rencontrera souvent les suivantes :

- transformation *isotherme* : $T = \text{cste}$,
- transformation *isobare* : $P = \text{cste}$,
- transformation *isochore* : $V = \text{cste}$,
- transformation *adiabatique* : $Q = 0$.

Une paroi qui ne permet pas le transfert thermique est dite *adiabatique* ; une paroi qui permet le transfert thermique est dite *diatherme* ou *diathermale*.

2.2 Transformations quasi-statiques

Une transformation *quasi-statique* est une transformation qui passe par une succession d'états d'équilibre très voisins les uns des autres. La transformation doit être assez lente pour que chaque équilibre ait le temps de s'établir.

On considère les deux transformations suivantes d'un gaz parfait contenu dans un cylindre fermé par un piston mobile. On dépose une masse m sur le piston soit progressivement (et lentement) en versant des petits cailloux (Cf. figure 36.3), soit d'un seul coup en utilisant un bloc de masse m (Cf. figure 36.4). Dans le premier cas, après la chute de chaque caillou, le système se met à l'équilibre, voisin du précédent : la transformation est quasi-statique. Dans le deuxième cas, il n'existe pas d'équilibre intermédiaire : la transformation n'est pas quasi-statique.

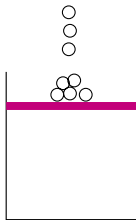


Figure 36.3 Ajout progressif et lent d'une masse sur un piston.

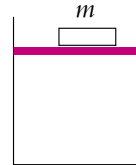


Figure 36.4 Ajout brutal d'une masse sur un piston.

Expérimentalement, on constate que l'équilibre mécanique est atteint beaucoup plus rapidement que l'équilibre thermique car les échanges thermiques sont lents (voir le cours de deuxième année sur la diffusion thermique). Sur l'exemple ci-dessous où $P_1 < P_2$, dès qu'on lâche la cloison diatherme, elle va quasi instantanément se déplacer vers la gauche jusqu'à ce que les pressions dans les deux compartiments soient égales puis les échanges thermiques vont avoir lieu plus lentement à travers la paroi. Cette transformation n'est pas quasi-statique.

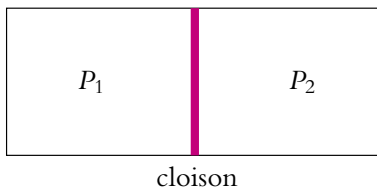


Figure 36.5 État initial : juste avant qu'on lâche la cloison.

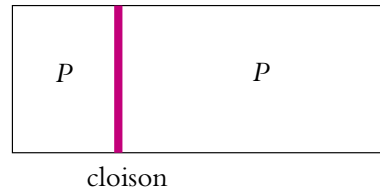


Figure 36.6 Dès qu'on lâche la cloison, elle se déplace pour assurer l'équilibre mécanique.

Par conséquent, on pourra souvent faire l'hypothèse que des transformations rapides sont adiabatiques, car les échanges thermiques n'ont pas eu le temps de se faire. Entre deux états d'équilibre voisins d'une transformation quasi-statique, on écrira le premier principe sous la forme :

$$dU = \delta W + \delta Q. \quad (36.5)$$

Il est très important de différencier les notations d et δ , qui mettent en valeur les propriétés différentes d'une fonction d'état comme U et de grandeurs comme W et Q .

- dU représente la variation d'énergie interne entre les deux états d'équilibres voisins (comme on écrit dE_c pour la variation élémentaire d'énergie cinétique en mécanique).
- δQ et δW représentent des petites quantités (et non des variations) de transfert thermique ou de travail reçus par le système.

C'est donc une faute grave d'écrire dQ et dW .

De même :

- ΔU représente la variation d'énergie interne entre les deux états d'équilibre initial et final ;
- Q et W représentent des quantités (et non des variations) de transfert thermique ou de travail reçus par le système.

C'est donc une faute grave d'écrire ΔQ et ΔW .

Si on intègre la relation (36.5) entre l'état initial et l'état final, on obtient l'expression du premier principe (36.1) :

$$\Delta U = W + Q$$

avec **ΔU variation d'énergie interne totale, Q transfert thermique total et W travail total reçus par le système.**

Ainsi, dans le cas de la transformation décrite sur la figure 36.3, entre deux états voisins $dU = \delta W + \delta Q$ et sur toute la transformation $\Delta U = W + Q$ tandis que dans le cas de celle décrite sur la figure 36.4, on ne peut pas utiliser la relation 36.5, et il faut écrire directement, les variables d'état n'étant définies que dans l'état initial et dans l'état final $\Delta U = W + Q$. Même si l'état initial est le même que dans l'autre exemple, l'état final, lui sera sûrement différent.

2.3 Transformations réversibles

On reviendra plus en détail sur la notion de transformation réversible dans le chapitre sur le deuxième principe. On signale cependant qu'une transformation est dite *réversible* si on peut inverser le sens d'évolution c'est-à-dire « revenir en arrière » et que le système repasse par exactement les mêmes états d'équilibre. On pourrait prendre l'image d'un film qu'on repasse à l'envers : si ce que l'on voit a une réalité physique, la transformation est réversible.

Les causes d'irréversibilité sont principalement les frottements, les gradients de température (le transfert thermique ne se fait spontanément que d'un corps chaud vers un corps froid) et les gradients de densité de particules (exemple d'une tache d'encre dans de l'eau).

Dans les deux exemples du paragraphe 2.2, la figure 36.3 correspond à une transformation réversible s'il n'y a pas de frottement, la figure 36.4 à une transformation irréversible. Une transformation réversible est quasi-statique mais une transformation quasi-statique n'est pas forcément réversible (par exemple, la transformation de la figure 36.3 s'il y a des frottements).

On peut citer deux autres exemples de transformations quasi-statiques et irréversibles, qui sont dues à des gradients de température ou de densité :

- On relie deux récipients contenant le même gaz par un tube très fin. Les pressions des deux récipients sont différentes mais la température est uniforme. En raison de la différence de pression, des particules quittent le récipient de plus forte pression pour aller, *via* le tube, dans l'autre récipient. À chaque instant, peu de particules changent de récipient à cause de la très faible section du tube. On peut considérer qu'il y a équilibre à chaque instant, la transformation est donc quasi-statique ; par contre, elle n'est pas réversible car les particules ne peuvent passer du récipient de faible pression à celui de forte pression.
- Le deuxième exemple est semblable mais cette fois avec deux solides à des températures différentes, reliés par un mauvais conducteur de chaleur. Il va s'établir un transfert thermique très lent. À chaque instant, on peut considérer que l'ensemble est quasiment à l'équilibre : la transformation est quasi-statique mais le transfert ne pouvant avoir lieu que dans un sens, elle n'est pas réversible.

3. Expression du travail

3.1 Système soumis à des forces de pression extérieures

Soit un système (S) contenu dans une enceinte déformable. On considère deux états de volumes voisins où un point M_1 de l'enceinte s'est déplacé en M_2 (Cf. figure 36.7) et on note $\vec{d\ell} = \overrightarrow{M_1 M_2}$ le petit déplacement du point. La pression extérieure est notée P_e . On considère la surface élémentaire dS , orientée par le vecteur $\vec{dS}$ autour du point M_1 qui se déplace avec lui (Cf. figure 36.8).

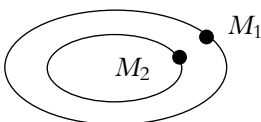


Figure 36.7 États initial et final de l'enceinte.

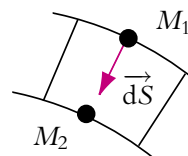


Figure 36.8 Déplacement de la surface dS .

La force exercée par l'extérieur sur l'élément de surface dS est $\vec{dF}$, son travail pendant la contraction de l'enceinte est $\delta^2 W = \vec{dF} \cdot \vec{d\ell} = P_e \vec{dS} \cdot \vec{d\ell}$ (On note le travail élémentaire $\delta^2 W$ car c'est le produit de deux infiniment petits). Or le produit scalaire $\vec{dS} \cdot \vec{d\ell}$ représente le volume balayé par la surface dS au cours de son déplacement. Pour toute l'enceinte, il suffit d'intégrer l'expression précédente : $\delta W = P_e \int_{M \in \text{enceinte}} \vec{dS}_M \cdot \vec{d\ell}_M$.

Or cette dernière intégrale, qui est positive dans le cas représenté, est égale à l'opposé de la variation de volume dV de l'enceinte d'où :

$$\delta W = -P_e dV \quad (36.6)$$

Si maintenant le système évolue de V_A (état E_A) à V_B (état E_B), on intègre l'équation 36.6 entre l'état E_A et l'état E_B . Le travail sera donc :

$$W = \int_{E_A}^{E_B} -P_e dV \quad (36.7)$$

Le travail qu'on vient de calculer est-il vraiment le travail reçu par (S) ? Non, pas exactement, c'est le travail reçu par le système (S') = enceinte + (S). Si on applique le premier principe à (S'), il faut tenir compte de la variation d'énergie interne de l'enceinte et de la variation de son énergie cinétique soit entre deux états voisins :

$$dU_S + dU_{\text{enceinte}} + dE_{\text{cinétique}} = -P_e dV + \delta Q \quad (36.8)$$

Or, dans la plupart des cas, l'enceinte est au repos dans chaque état ($dE_{\text{cinétique}} = 0$) et on néglige sa variation d'énergie interne car sa capacité thermique est très faible devant celle de (S) ($dU_{\text{enceinte}} = 0$) d'où d'après (36.8) :

$$dU_S = -P_e dV + \delta Q. \quad (36.9)$$

et l'on peut considérer que le travail calculé (36.6) est celui reçu par le gaz :

$$\delta W = -P_e dV$$

Dans le cas d'une transformation quasi-statique, la pression P du gaz est en permanence égale à la pression extérieure P_e , et

$$\delta W = -P dV. \quad (36.10)$$



Dans la pression extérieure interviennent toutes les forces extérieures, par exemple le poids d'un piston (Cf. figure 36.9). La masse du piston est m , la section du cylindre S et la pression de l'air P_a . La pression extérieure s'exerçant sur le système constitué du piston et du gaz du cylindre est $P_e = P_a + \frac{mg}{S}$.

Exceptions importantes aux formules (36.6) et (36.7) :

Peut-on toujours confondre le volume balayé et la variation de volume ?

La réponse est non. Il y a deux contre-exemples importants : celui où existe une enceinte vide au départ et celui d'un fluide en écoulement (voir paragraphe 5 du chapitre suivant). Il faut alors revenir à l'expression du travail comme le produit de la pression extérieure P_e par le volume balayé ou à la définition mécanique du travail d'une force.

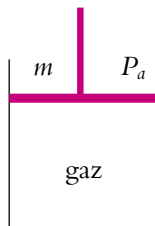


Figure 36.9 Prise en compte de la masse du piston dans le calcul de la pression extérieure.

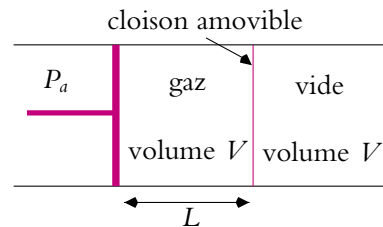


Figure 36.10 Exception aux formules (36.6) et (36.7).

On traite ici le cas d'une enceinte vide initialement (Cf. figure 36.10), l'autre cas étant traité dans le paragraphe indiqué précédemment. Le piston est mobile et la cloison entre les deux récipients de volume V est amovible. On enlève la cloison et on considère le cas d'une transformation isotherme.

Le système à considérer est le système (gaz + enceinte + piston). À l'état initial, le piston est au repos donc la pression du gaz est P_a , son volume V et sa température T . À l'état final, la température est T , le piston étant de nouveau immobile la pression est encore P_a , le volume est donc V et le gaz occupe toute la partie droite. Le piston s'est déplacé d'une longueur L et le travail de la force de pression est donc $W = FL = P_a SL$ (S étant la section de la cloison) soit $W = P_a V$ alors que la variation du volume occupé par le gaz ici est nulle.

En résumé : Dans le cas où l'on néglige l'intervention de l'enceinte dans le bilan d'énergie, le travail reçu par le fluide peut s'écrire $\delta W = P_e V_{\text{balayé}}$. Dans la plupart des cas, c'est-à-dire s'il n'y a pas d'enceinte vide au départ ou si on n'a pas un fluide en écoulement, la variation de volume est égale au volume balayé et on peut écrire $\delta W = -P_e dV$. Enfin, dans le cas d'une transformation quasi-statique, la pression P du gaz est égale à la pression extérieure et $\delta W = -P dV$.

3.2 Représentation graphique

On considère la représentation graphique de la transformation dans le diagramme de Watt donnant la pression P en fonction du volume V . Si on peut représenter la trans-

formation par une courbe, c'est qu'elle est constituée d'une suite d'états d'équilibre donc qu'elle est quasi-statique.

On représente une évolution de l'état A à l'état B (Cf. figure 36.11). La transformation étant quasi-statique, le travail entre A et B s'écrit $W_{A \rightarrow B} = \int_A^B -P dV$ ce qui représente l'opposé de l'aire sous la courbe AB . On constate donc que, dans le cas représenté, le travail reçu par le système est négatif si $V_B > V_A$ et positif dans le cas contraire.

On s'intéresse maintenant au cas d'une transformation cyclique (Cf. figure 36.12). Par le chemin (1) allant de A à B , le système reçoit W_1 qui est l'opposé de l'aire S_1 sous la courbe (1), puis pour passer de B à A par le chemin (2), le système reçoit le travail W_2 égal à l'aire S_2 sous la courbe (2). Sur le cycle, le travail reçu est $W = W_1 + W_2$ donc égal à $S_2 - S_1$ qui est l'opposé de l'aire du cycle.

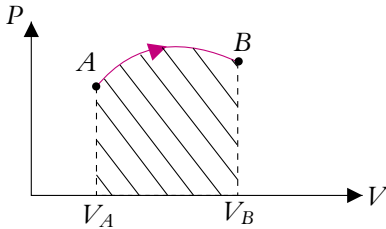


Figure 36.11 Interprétation graphique du travail des forces de pression.

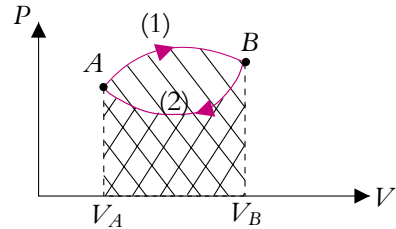


Figure 36.12 Cycles moteur ou récepteur et travail des forces de pression.

Dans le cas représenté le travail W est donc négatif. On dit que le cycle est *moteur*.

Si on avait considéré un cycle tournant dans le sens horaire, on aurait eu $W_1 = S_1$, $W_2 = -S_2$ et $W = S_1 - S_2 > 0$, le travail est donc égal à l'aire du cycle et le cycle est récepteur.

Si un cycle « tourne » dans le sens horaire, il s'agit d'un cycle moteur, s'il « tourne » dans le sens trigonométrique, il s'agit d'un cycle récepteur.

On remarque que le travail n'est pas nul sur un cycle alors que la variation de l'énergie interne $\Delta U = U_B - U_A$ est nulle. **Ceci est dû au fait que W n'est pas la variation d'une fonction d'état.**

► **Remarque :** on doit distinguer le *diagramme de Watt* (Pression en fonction du volume) et le *diagramme de Clapeyron* (Pression en fonction du volume massique, c'est-à-dire pour un kilogramme de fluide). Cependant, il arrive fréquemment que l'on utilise l'appellation diagramme de Clapeyron pour désigner le diagramme de Watt.

3.3 Cas d'une transformation isochore

Lorsque la transformation est isochore, il n'y a pas de variation de volume, le travail des forces de pression est nul d'après ce qui précède donc :

$$\Delta U = U_2 - U_1 = Q + W_{\text{autre}} \quad (36.11)$$

où W_{autre} est le travail des forces autres que les forces de pression.

Dans le cas fréquent où le travail des forces autres que les forces de pression est nul, $\Delta U = Q$ au cours d'une transformation isochore.

3.4 Travail électrique

Il arrive souvent qu'on rencontre des systèmes recevant de l'énergie électrique (pile, eau dans calorimètre...).

Pour être cohérent avec la convention d'orientation de Q et W , on choisit **une convention récepteur** (voir cours d'électricité).

Le travail reçu par le système (S) soumis à une tension $u(t)$ et parcouru par une intensité $i(t)$, entre deux instants t_1 et t_2 :

$$W = \int_{t_1}^{t_2} u(t)i(t)dt \quad (36.12)$$

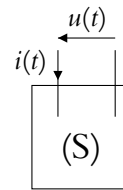


Figure 36.13 Orientation récepteur pour un système électrique.

4. Enthalpie

4.1 Transformation monobare et transformation isobare

On appelle transformation *monobare* une transformation pour laquelle la pression extérieure est constante.

Au cours de la transformation, la pression du système peut varier. Une transformation isobare est un cas particulier d'une transformation monobare quasistatique au cours de laquelle la pression du système reste constante. De nombreuses réactions chimiques ont lieu à pression extérieure constante et sont donc monobares, d'où l'utilité de cette étude. On considère donc un système subissant une évolution monobare, sous la pression extérieure constante P_e entre les deux états suivants :

$$\begin{array}{c} \text{État initial (1)} \end{array} \left| \begin{array}{c} P_1 \\ V_1 \\ T_1 \end{array} \right. \quad \parallel \quad \begin{array}{c} \text{État final (2)} \end{array} \left| \begin{array}{c} P_2 \\ V_2 \\ T_2 \end{array} \right.$$

On aimerait trouver une fonction dont la variation soit égale au transfert thermique à l'image de la relation (36.11) lorsqu'il n'y a que des forces de pression. On applique le premier principe au système entre l'état (1) et l'état (2) : $\Delta U = U_2 - U_1 = W + Q$. Or la pression extérieure P_e est constante, donc le travail des forces de pression est

$$W = -P_e(V_2 - V_1) \text{ soit}$$

$$(U_2 + P_e V_2) - (U_1 + P_e V_1) = Q$$

On a obtenu ce que l'on cherchait en définissant la fonction $H^* = U + P_e V$. Ainsi :

$$\Delta H^* = Q$$

4.2 Définition de l'enthalpie

Cependant, H^* n'est pas une fonction d'état car elle dépend de la pression extérieure P_e . On envisage maintenant le cas où l'état initial et l'état final sont des états d'équilibre, c'est-à-dire le cas où $P_1 = P_e$ et $P_2 = P_e$. On peut alors définir une fonction d'état qui coïncide avec H^* à l'état final et à l'état initial lors d'une transformation monobare telle que l'état initial et l'état final soient des états d'équilibre. Cette fonction, appelée *enthalpie* et notée H , est définie par :

$$H = U + PV \quad (36.13)$$

où U est l'énergie interne, P la pression et V le volume.

Comme on l'a remarqué au paragraphe 3.1, le produit PV est homogène à une énergie et donc à U . L'unité de l'enthalpie est donc le joule.

Comme pour l'énergie interne U , la variation d'enthalpie H ne dépend que des états extrêmes et non du chemin suivi, c'est une fonction d'état.

Dans le cas de la transformation monobare précédente : $P_1 = P_2 = P_e$ et donc :

$$\Delta H = H_2 - H_1 = Q \quad (36.14)$$

Cette dernière relation est vérifiée pour une transformation isobare qui est un cas particulier de transformation monobare.

Attention la relation (36.14) n'est valable que dans les cas où il s'exerce sur le corps uniquement des forces de pression, puisque on n'a tenu compte que de ce type de forces dans le calcul du travail. Dans le cas contraire, il faut ajouter le travail autre que celui des forces de pression (par exemple électrique), appelé W_{autre} . L'expression (36.14) devient :

$$\Delta H = H_2 - H_1 = Q + W_{\text{autre}} \quad (36.15)$$

4.3 Capacité thermique à pression constante

a) Gaz parfait monoatomique

On prend comme système un gaz parfait monoatomique. Dans ce cas, on peut écrire les relations suivantes :

$$dU = \frac{3}{2} nRdT$$

et

$$dH = dU + d(PV) = dU + d(nRT) = dU + nRdT = \frac{5}{2}nRdT$$

On définit la capacité thermique à pression constante pour un gaz parfait monoatomique par :

$$dH = C_P dT \quad \text{et} \quad C_P = \frac{5}{2}nR \quad (36.16)$$

b) Gaz parfait quelconque

On retrouve une définition similaire à celle de C_V avec la fonction U . Pour un gaz parfait quelconque, on peut établir une relation entre C_P et C_V :

$$dU = C_V dT$$

et

$$dH = dU + d(PV) = dU + d(nRT) = dU + nRdT = (C_V + nR)dT$$

d'où

$$C_P - C_V = nR \quad (36.17)$$

Cette relation, appelée *relation de Mayer* pour le gaz parfait, montre que $C_P > C_V$. Dans le cas d'un gaz quelconque, on définit le rapport des capacités thermiques

$$\gamma = \frac{C_P}{C_V} \quad (36.18)$$

Comme C_P est supérieur à C_V (Cf. 36.17), γ est toujours plus grand que 1. En général dans un problème, on donne γ . Il faut donc savoir exprimer C_P et C_V en fonction de γ . Pour cela on utilise les relations (36.17) et (36.18) dont on déduit pour un gaz parfait :

$$C_V = \frac{nR}{\gamma - 1} \quad \text{et} \quad C_P = \frac{nR\gamma}{\gamma - 1} \quad (36.19)$$

L'unité de C_V et C_P est le joule par kelvin.

On peut définir les capacités thermiques molaires C_{Vm} et C_{Pm} en divisant par le nombre de moles n soit $C_{Vm} = \frac{R}{\gamma - 1}$ et $C_{Pm} = \frac{\gamma R}{\gamma - 1}$, ou massiques en divisant par la masse soit $c_V = \frac{R}{M(\gamma - 1)}$ et $c_P = \frac{R\gamma}{M(\gamma - 1)}$, où M est la masse molaire.

Le fait que l'enthalpie d'un gaz parfait ne dépende que de la température est appelée *seconde loi de Joule*.

$$\Delta H = C_P \Delta T \quad (36.20)$$

On rappelle que la première loi de Joule (35.17) correspond au fait que l'énergie interne d'un gaz parfait ne dépend elle aussi que de la température.



Les formules (36.19) ne sont valables que pour un gaz parfait.

c) Cas général

On appelle C_p *capacité thermique à pression constante* car, de manière générale, H dépend de deux variables, comme U . On choisit en général P et T pour décrire l'enthalpie H . Pour tout système, on peut alors définir la capacité thermique à pression constante par :

$$C_p = \left(\frac{\partial H}{\partial T} \right)_p \quad (36.21)$$

Il s'agit donc de la variation d'enthalpie lorsque la température varie d'un kelvin sous une pression constante.

Comme cela a été signalé précédemment, l'unité de C_p est la même que celle de C_V : J.K^{-1} . En résumé, les capacités thermiques s'expriment en J.K^{-1} ; les capacités molaires en $\text{J.K}^{-1}.\text{mol}^{-1}$ et les capacités massiques en $\text{J.K}^{-1}.\text{kg}^{-1}$.

► **Remarque :** Il n'y a pas de convention de notation pour les capacités thermiques massiques ou molaires. On aura intérêt à faire attention aux unités dans les problèmes pour savoir de quelle quantité il s'agit.

d) Cas des solides et des liquides peu compressibles

Pour les solides et les liquides peu compressibles, on peut négliger les variations de volume et donc celles du produit PV . Dans ce cas, les variations d'enthalpie sont pratiquement égales aux variations d'énergie interne puisque : $dH = dU + d(PV)$ et que $d(PV) \ll dU$. On peut s'en convaincre sur un exemple :

Exemple

À 0°C , sous une pression de $1 \text{ bar} = 10^5 \text{ Pa}$, le volume d'un kilogramme d'eau liquide est environ $1 \text{ L} = 10^{-3} \text{ m}^3$, alors qu'à 600°C sous une pression de $200 \text{ bar} = 200.10^5 \text{ Pa}$, le volume de cette même quantité est $3 \text{ L} = 3.10^{-3} \text{ m}^3$. Sachant que la capacité thermique de l'eau liquide est quasiment indépendante de la température et égale à $4,18 \text{ kJ.kg}^{-1}$, on déduit qu'entre 0°C et 600°C , pour un kilogramme d'eau liquide :

$$\Delta U = 4,18.10^3 \times 600 = 2,51.10^6 \text{ J} \quad \text{et} \quad \Delta H = \Delta U + (PV)_{600^\circ\text{C}} - (PV)_{0^\circ\text{C}}$$

soit :

$$\Delta H = \Delta U + 6.10^4 - 10^2 = 2,57.10^6 \text{ J}$$

Ainsi, sur une plage très vaste allant de 0°C à 600°C , si on confond ΔU et ΔH on commet une erreur relative de 2 %.

Conclusion : dans la pratique pour un solide ou un liquide on confond enthalpie H et énergie interne U et on donne une capacité thermique C sans préciser si elle est à volume constant ou à pression constante.

Alors pour une transformation élémentaire :

$$dH = dU = C dT \quad (36.22)$$

et entre l'état initial et l'état final :

$$\Delta H = \Delta U = C \Delta T \quad (36.23)$$

si on considère C indépendante de la température dans le domaine de températures considéré.

Applications du premier principe

37

1. Calorimétrie

Un *calorimètre* est un récipient composé en général d'une paroi extérieure et d'une cuve intérieure, fermé par un couvercle permettant d'introduire un agitateur, un thermomètre, une résistance chauffante. La cuve intérieure étant séparée de la paroi extérieure par de l'air, le système est relativement bien isolé et on peut négliger sur le temps d'une expérience en travaux pratiques les échanges thermiques avec l'extérieur.

En revanche, au paragraphe 3.1 du chapitre précédent, on a écrit qu'on pouvait en général négliger la variation d'énergie interne d'une enceinte car sa capacité thermique était petite devant celle d'un gaz. Ce n'est pas le cas pour un calorimètre : on doit donc tenir compte de sa capacité thermique. Par habitude, au lieu de donner la capacité thermique du calorimètre, on donne la masse d'eau qui aurait la même capacité thermique et qu'on appelle *valeur en eau du calorimètre*. Ainsi, un calorimètre ayant une valeur en eau $\mu = 20 \text{ g}$ a une capacité thermique $C = 4,18 \cdot 10^3 \times 20 \cdot 10^{-3} = 83,6 \text{ J.K}^{-1}$.

La valeur en eau du calorimètre tient compte de la capacité thermique de tous les instruments du calorimètre.

Quelle fonction d'état utiliser ? Les expériences dans un calorimètre se font à pression extérieure constante, le système étant en contact avec l'atmosphère. On utilisera donc l'enthalpie et la relation (36.15).

Toutes les températures t seront exprimées en degré Celsius.

1.1 Méthode des mélanges ; détermination de la capacité thermique d'un solide

On considère l'expérience suivante : un calorimètre de valeur en eau μ parfaitement isolé contient une masse m d'eau liquide à la température t_0 . On plonge un corps de capacité massique c_P , de masse m_C et de température t_1 . On attend le retour à l'équilibre et on note t_F la température finale. Peut-on déterminer la capacité thermique massique c_P du corps ?

Soit c_E la capacité thermique massique de l'eau liquide.

Système : calorimètre et instruments + eau + corps solide

$$\begin{array}{c|c|c} \text{État initial} & \begin{array}{l} \text{Eau + calorimètre : } t_0 \\ \text{Corps solide : } t_1 \end{array} & \text{État final} \begin{array}{l} \text{Eau + calorimètre : } t_F \\ \text{Corps solide : } t_F \end{array} \end{array}$$

Application du premier principe pour une évolution monobare :

$$\Delta H = W_{\text{autre}} + Q$$

L'enthalpie H est une fonction extensive donc :

$$\Delta H = \Delta H_{\text{eau}} + \Delta H_{\text{calo}} + \Delta H_{\text{corps}}$$

Il n'y pas d'autre force que les forces de pression, donc $W_{\text{autre}} = 0$ et comme le calorimètre est isolé, $Q = 0$. Ensuite il suffit de calculer ΔH pour les trois corps en utilisant les résultats du paragraphe 4.3.d) du chapitre précédent et les conditions initiales et finales.

$$\Delta H_{\text{eau}} = mc_E(t_F - t_0), \quad \Delta H_{\text{calo}} = \mu c_E(t_F - t_0), \quad \Delta H_{\text{corps}} = m_C c_P(t_F - t_1)$$

soit, avec $\Delta H = 0$:

$$c_P = -\frac{(m + \mu)c_E(t_F - t_0)}{m_C(t_F - t_1)} \quad (37.1)$$

► **Remarque :** Il n'y a pas à se poser de question sur le sens des échanges d'énergie. Il suffit de définir l'état initial, l'état final et d'appliquer le premier principe avec les conventions précédemment définies. La température finale t_f est comprise entre t_1 et t_0 .

1.2 Méthode électrique

On cherche à mesurer la valeur en eau d'un calorimètre, valeur indispensable pour toute autre mesure. Le calorimètre dont on cherche la valeur en eau μ est parfaitement isolé et contient une masse m d'eau liquide à la température t_0 . On plonge dans l'eau liquide une résistance chauffante R , qui est parcourue par un courant I pendant l'intervalle de temps τ .

On note c_E la capacité thermique massique de l'eau liquide. On néglige la capacité thermique de la résistance. Cette dernière reçoit du générateur un travail $W_{\text{elec}} = RI^2\tau$.

Système : calorimètre et instruments + eau + résistance

État initial | Eau + calorimètre : t_0 || État final | Eau + calorimètre : t_F

Application du premier principe pour une évolution monobare :

$$\text{avec} \quad \Delta H = W_{\text{autre}} + Q$$

$$\Delta H = \Delta H_{\text{eau}} + \Delta H_{\text{calo}}$$

Or ici $W_{\text{autre}} = W_{\text{elec}}$ et, comme le calorimètre est isolé, $Q = 0$. Ensuite il suffit de calculer ΔH comme au paragraphe 1.1 :

$$\Delta H_{\text{eau}} = mc_E(t_F - t_0), \quad \Delta H_{\text{calo}} = \mu c_E(t_F - t_0)$$

soit avec $\Delta H = W_{\text{elec}} = Ri^2\tau$

$$\mu = \frac{Ri^2\tau}{c_E(t_F - t_0)} - m \quad (37.2)$$

1.3 Cas d'un changement d'état

On considère l'expérience suivante : un calorimètre de valeur en eau μ , parfaitement isolé, contient une masse m d'eau liquide à la température t_0 . On met dans le calorimètre une masse m_g de glace à température t_1 . Quel est l'état final, la pression étant égale à 1 bar ?

Le chapitre 40 sera consacré au changement de phase mais on admet ici que la fusion de la glace s'accompagne d'une variation d'enthalpie. Intuitivement on sait qu'il faut apporter de l'énergie pour faire fondre de la glace.

On définit l'*enthalpie massique de fusion* (attention, c'est en fait une variation d'enthalpie) comme étant la variation d'enthalpie d'un kilogramme de solide passant à l'état liquide à pression et température constantes.

On la note l_F et c'est une grandeur positive.

Ainsi, pour la fusion d'une masse m_g de glace, on a $\Delta H = m_g l_F$ et, pour la solidification d'une masse m_g de liquide, on a $\Delta H = -m_g l_F$. Il faut donc faire attention au sens du changement d'état.

De même, on peut définir une *enthalpie massique de vaporisation*, comme la variation d'enthalpie d'un kilogramme de liquide passant à l'état vapeur à pression et température constante. On la note l_V et c'est aussi une grandeur positive. Pour la vaporisation

d'une masse m_l de liquide on a $\Delta H = m_l l_V$ et, pour la condensation d'une masse m_l de vapeur, on a $\Delta H = -m_l l_V$.

Dans le cas précis de l'expérience, on ne sait pas exactement par quelles étapes passe le système lors de son évolution.

On utilise le fait que l'enthalpie est une fonction d'état et que sa variation est indépendante du chemin suivi et dépend seulement de l'état initial et de l'état final. On imagine une évolution partant de l'état initial connu et aboutissant à l'état final probable. Ici trois états finaux sont possibles :

1. eau entièrement liquide à température supérieure à 0°C ,
2. mélange liquide glace à 0°C ,
3. eau entièrement solide à température inférieure à 0°C .

On fait une hypothèse sur l'état final et on vérifie après calculs que le résultat est en accord avec l'hypothèse. Si, par exemple, il s'agit de l'hypothèse 1, l'inconnue est la température finale, mais il faudra vérifier qu'elle est supérieure à 0°C . S'il s'agit de l'hypothèse 2, la masse finale de glace est inconnue et il faudra vérifier qu'elle n'est pas supérieure à la masse totale d'eau. On choisit l'hypothèse 1 avec comme données :

- la capacité thermique massique de la glace : c_g ,
- la capacité thermique massique de l'eau liquide : c_l ,
- l'enthalpie massique de fusion de la glace sous 1 bar à 0°C : l_f .

Système : calorimètre et instruments + eau liquide + glace

État initial	Eau liquide : m, t_0 Calorimètre : μ, t_0 Glace : m_g, t_1	État final	Eau liquide : $m + m_g, t_F$ Calorimètre : t_F
--------------	---	------------	---

Application du premier principe pour une évolution monobare :

$$\Delta H = W_{\text{autre}} + Q$$

avec

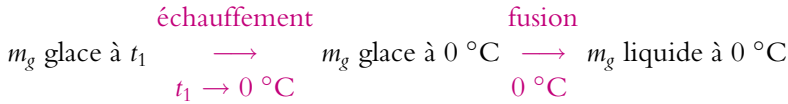
$$\Delta H = \Delta H_{\text{eau}} + \Delta H_{\text{calo}} + \Delta H_{\text{mg}}$$

Or ici, comme au paragraphe 1.1, il n'y pas d'autre force que les forces de pression donc $W_{\text{autre}} = 0$ et comme le calorimètre est isolé, $Q = 0$. Les calculs de ΔH_m (masse m d'eau initialement liquide) et ΔH_{calo} s'effectuent comme au paragraphe 1.1. :

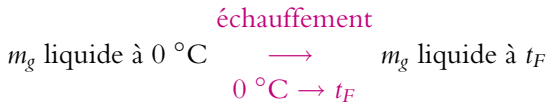
$$\Delta H_m = mc_l(t_F - t_0) \quad \text{et} \quad \Delta H_{\text{calo}} = \mu c_l(t_F - t_0)$$

Pour le calcul de ΔH_{mg} (masse m d'eau initialement glace), il faut imaginer une évolution partant de l'état initial connu, allant à l'état final supposé, sachant qu'on donne

l'enthalpie de fusion à 0°C . On imagine l'évolution suivante :



puis



On peut alors calculer $\Delta H_{m_g} = m_g c_g(0 - t_1) + m_g l_f + m_g c_l(t_F - 0)$.

Avec les expressions précédentes, on peut alors déterminer t_F et vérifier l'hypothèse. Si, disposant de données numériques, on trouve que $t_F < 0^\circ\text{C}$, c'est que toute la glace n'est pas fondue et qu'il faut choisir une autre hypothèse.

2. Transformations d'un gaz parfait

On passe en revue la plupart des transformations rencontrées et on calcule dans chaque cas le transfert thermique reçu Q , le travail reçu W , la variation d'énergie interne ΔU et la variation d'enthalpie ΔH . Le système considéré est un gaz parfait formé de n moles.

La démarche suivie sera toujours la même : on calcule la variation d'énergie interne, la variation d'enthalpie et le travail reçu puis on en déduit le transfert thermique reçu. Bien sûr, si on connaît la valeur de Q (comme dans le cas d'une transformation adiabatique), on calculera W à partir de l'expression du premier principe. On considère que le seul travail à prendre en compte est celui des forces de pression.

2.1 Transformation quasi-statique isotherme d'un gaz parfait

$$\text{État initial (1)} \left\{ \begin{array}{l} V_1 \\ P_1 \\ T_1 = T \end{array} \right. \quad \parallel \quad \text{État final (2)} \left\{ \begin{array}{l} V_2 \\ P_2 \\ T_2 = T \end{array} \right.$$

La transformation est isotherme : la température T du système reste constante.

Entre deux états intermédiaires, $\delta W = -PdV = -\frac{nRT}{V}$, soit entre les deux états extrêmes :

$$W = -nRT \int_{V_1}^{V_2} \frac{dV}{V} = -nRT \ln \left(\frac{V_2}{V_1} \right) = nRT \ln \left(\frac{P_2}{P_1} \right) \quad (37.3)$$

D'autre part, en raison des deux lois de Joule :

$$\Delta U = C_V \Delta T = 0 \quad \text{et} \quad \Delta H = C_P \Delta T = 0 \quad (37.4)$$

et comme $\Delta U = W + Q$

$$Q = -W = nRT \ln \left(\frac{V_2}{V_1} \right) = -nRT \ln \left(\frac{P_2}{P_1} \right)$$

► **Remarque :** Comme on l'a déjà remarqué, il peut y avoir transfert thermique sans variation de température. Dans le cas d'une compression isotherme ($P_2 > P_1$ et $V_2 < V_1$), alors $W > 0$ et $Q < 0$. Il faut fournir effectivement du travail au système qui fournit du transfert thermique à l'extérieur. C'est le contraire dans le cas d'une détente isotherme.

2.2 Transformation quasi-statique isochore d'un gaz parfait

$$\begin{array}{c|c|c} \text{État initial (1)} & \left| \begin{array}{l} V_1 = V \\ P_1 \\ T_1 \end{array} \right. & \parallel \parallel \parallel \text{État final (2)} \left| \begin{array}{l} V_2 = V \\ P_2 \\ T_2 \end{array} \right. \end{array}$$

Entre deux états intermédiaires $\delta W = -PdV = 0$ soit :

$$W = 0$$

► **Remarque :** Cette dernière relation est valable pour tout autre système qu'un gaz parfait.

Pour un gaz parfait, il vient alors :

$$\Delta U = C_V(T_2 - T_1) = Q \quad (37.5)$$

$$\Delta H = C_P(T_2 - T_1) \quad (37.6)$$

2.3 Transformation quasi-statique isobare d'un gaz parfait

$$\begin{array}{c|c|c} \text{État initial (1)} & \left| \begin{array}{l} V_1 \\ P_1 = P \\ T_1 \end{array} \right. & \parallel \parallel \parallel \text{État final (2)} \left| \begin{array}{l} V_2 \\ P_2 = P \\ T_2 \end{array} \right. \end{array}$$

Entre deux états intermédiaires, $\delta W = -PdV$ soit

$$W = -P(V_2 - V_1)$$

► **Remarque :** Cette dernière relation est valable pour tout autre système qu'un gaz parfait.

Pour un gaz parfait, il vient alors :

$$\Delta U = C_V(T_2 - T_1) \quad (37.7)$$

$$\Delta H = C_P(T_2 - T_1) = Q \quad (37.8)$$

2.4 Transformation quasi-statique adiabatique d'un gaz parfait

$$\begin{array}{c} \text{État initial (1)} \left| \begin{array}{c} V_1 \\ P_1 \\ T_1 \end{array} \right| \quad \quad \quad \text{État final (2)} \left| \begin{array}{c} V_2 \\ P_2 \\ T_2 \end{array} \right| \\ Q = 0 \end{array}$$

donc

$$W = \Delta U = \frac{nR}{\gamma - 1}(T_2 - T_1) = \frac{P_2 V_2 - P_1 V_1}{\gamma - 1}$$

et

$$\Delta H = C_P(T_2 - T_1) = \frac{nR\gamma}{\gamma - 1}(T_2 - T_1) = \frac{\gamma}{\gamma - 1}(P_2 V_2 - P_1 V_1)$$

D'autre part, entre deux états intermédiaires, $\delta Q = 0$ soit $dU = \delta W$ et $C_V dT = -P dV \Leftrightarrow \frac{nR}{\gamma - 1} dT = -\frac{nRT}{V} dV$. En intégrant entre les états (1) et (2) en considérant γ indépendant de la température :

$$\int_{T_1}^{T_2} \frac{dT}{T} = -(\gamma - 1) \int_{V_1}^{V_2} \frac{dV}{V} \Leftrightarrow \ln\left(\frac{T_2}{T_1}\right) = -(\gamma - 1) \ln\left(\frac{V_2}{V_1}\right) \Leftrightarrow T_2 V_2^{\gamma-1} = T_1 V_1^{\gamma-1}$$

L'expression obtenue constitue l'une des équations de Laplace qui sont au nombre de trois. Il suffit d'utiliser l'équation du gaz parfait pour obtenir les deux autres.

Équations de Laplace :

$$\begin{aligned} T_2 V_2^{\gamma-1} &= T_1 V_1^{\gamma-1} \\ P_2 V_2^\gamma &= P_1 V_1^\gamma \\ T_2^\gamma P_2^{1-\gamma} &= T_1^\gamma P_1^{1-\gamma} \end{aligned} \quad (37.9)$$



Les équations de Laplace ne sont valables que pour le gaz parfait, dans le cas d'une transformation quasi-statique et adiabatique d'un système fermé, et si γ est indépendant de la température.

2.5 Comparaison des pentes d'une isotherme et d'une adiabatique dans le diagramme de Watt (P, V)

Il s'agit de calculer la dérivée de la pression P par rapport au volume V dans le cas d'une évolution isotherme et dans le cas d'une évolution adiabatique, en un même point (P, V) . On utilise les résultats précédents des paragraphes 2.1 et 2.4.

- Pour une isotherme $PV = nRT = \text{cste}$ donc la pente de l'isotherme est :

$$\left(\frac{\partial P}{\partial V}\right)_T = -\frac{nRT}{V^2} = -\frac{P}{V}$$

- Pour une transformation adiabatique quasi-statique d'un gaz parfait : $PV^\gamma = \text{cste}'$ donc la pente de l'adiabatique est :

$$\left(\frac{\partial P}{\partial V}\right)_{\text{adiabatique}} = -\frac{\gamma \text{cste}'}{V^{\gamma+1}} = -\gamma \frac{P}{V}$$

Or $\gamma > 1$ donc l'adiabatique est plus pentue que l'isotherme.

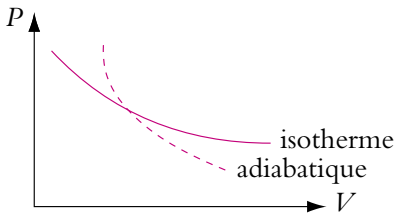


Figure 37.1 Comparaison des pentes d'une isotherme et d'une adiabatique.

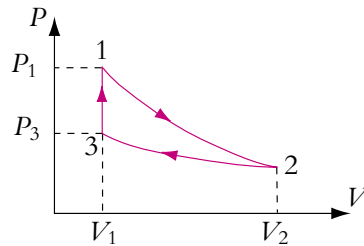


Figure 37.2 Cycle comportant une adiabatique, une isotherme et une isochore.

Ainsi on obtient la figure 37.1 quand une adiabatique et une isotherme se croisent en un point. L'adiabatique est en trait pointillé et l'isotherme en trait plein.

On peut appliquer ceci pour former un cycle constitué d'une transformation adiabatique, d'une transformation isotherme et d'une transformation isochore (Cf. figure 37.2). On calcule alors le travail sur le cycle : $W_{\text{cycle}} = W_{12} + W_{23} + W_{31}$ soit $W_{\text{cycle}} = \frac{P_2 V_2 - P_1 V_1}{\gamma - 1} + nRT_2 \ln\left(\frac{V_2}{V_1}\right) + 0$. Sur le cycle $\Delta U = U_1 - U_1 = 0$ donc $Q = -W$

2.6 Cycle de Carnot réversible d'un gaz parfait

Un cycle de Carnot est constitué de deux transformations isothermes et de deux transformations adiabatiques réversibles. Les transferts thermiques ont donc lieu avec uniquement deux sources à températures constantes. On parle alors d'un *cycle ditherme*. On appellera T_1 la température de la source chaude et T_2 celle de la source froide.

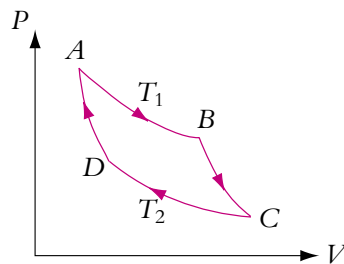


Figure 37.3 Cycle de Carnot.

- AB : transformation isotherme à la température T_1 ,
- BC : transformation adiabatique,
- CD : transformation isotherme à la température T_2 ,
- DA : transformation adiabatique.

Système : le gaz décrivant le cycle

On veut déterminer le travail total W et trouver une relation entre les transferts thermiques Q_1 et Q_2 reçus par le gaz entre A et B et entre C et D .

On applique le premier principe au gaz sur le cycle : $\Delta U_{\text{cycle}} = 0 = W + Q_1 + Q_2$ puisque les transformations BC et DA sont adiabatiques (Q_{BC} et Q_{DA} sont nuls).

► **Remarque :** la variation d'énergie interne est nulle sur le cycle car U est une fonction d'état.

Pour calculer le travail total W reçu par le fluide au cours du cycle, on a donc intérêt à calculer Q_1 et Q_2 (deux calculs) au lieu des travaux sur les quatre transformations.

Les transformations AB et CD sont isothermes donc, avec les résultats (37.3) :

$$\begin{aligned} Q_1 = Q_{AB} = -W_{AB} &= nRT_1 \ln \left(\frac{P_A}{P_B} \right) \\ Q_2 = Q_{CD} = -W_{CD} &= nRT_2 \ln \left(\frac{P_C}{P_D} \right) \end{aligned} \quad (37.10)$$

Les transformations BC et DA sont adiabatiques réversibles, on peut donc appliquer les lois de Laplace (37.9) au gaz parfait :

$$\begin{aligned} T_2^\gamma P_D^{1-\gamma} &= T_1^\gamma P_A^{1-\gamma} \\ T_2^\gamma P_C^{1-\gamma} &= T_1^\gamma P_B^{1-\gamma} \end{aligned} \quad (37.11)$$

On forme maintenant l'expression $\frac{Q_1}{T_1} + \frac{Q_2}{T_2}$. Les expressions (37.10) donnent :

$$\frac{Q_1}{T_1} + \frac{Q_2}{T_2} = nR \ln \left(\frac{P_A P_C}{P_B P_D} \right)$$

Si on utilise les relations (37.11), l'expression précédente devient :

$$\ln \left(\frac{T_1 T_2}{T_1 T_2} \right)^{\frac{\gamma}{1-\gamma}} = 0$$

Ainsi la relation entre les transferts thermiques est :

$$\frac{Q_1}{T_1} + \frac{Q_2}{T_2} = 0 \quad (37.12)$$

Le cycle qui est dessiné est **un cycle moteur** (il tourne dans le sens des aiguilles d'une montre). On remarque que le système (le gaz qui décrit le cycle) reçoit effectivement du transfert thermique de la source chaude ($Q_1 > 0$) et en cède à la source froide ($Q_2 < 0$). Pour un moteur de voiture, Q_1 provient de l'explosion de l'essence et la source froide est l'air extérieur. On note cependant qu'un moteur de voiture ne fonctionne pas selon un cycle de Carnot. Les signes des transferts thermiques sont toujours les mêmes pour un cycle moteur comme cela sera vu dans le chapitre sur les machines thermiques. La relation (37.12) permet d'écrire $Q_1 = -\frac{T_1}{T_2} Q_2$. On en déduit que $|Q_1| > |Q_2|$ et que **le système fonctionnant en moteur reçoit plus d'énergie de la source chaude qu'il n'en cède à la source froide**. D'après le premier principe, la différence correspond au travail récupérable $W = -(Q_1 + Q_2)$. On peut définir le rendement du moteur comme étant le travail récupérable divisé par le transfert thermique reçu :

$$\eta = \frac{|W|}{|Q_1|} = 1 + \frac{Q_2}{Q_1} = 1 - \frac{T_2}{T_1} \quad (37.13)$$

Si le cycle tournait dans le sens trigonométrique, on aurait un cycle récepteur (réfrigérateur, pompe à chaleur, climatiseur...). Dans ce cas, le travail W serait positif, le transfert thermique Q_1 négatif et le transfert thermique Q_2 positif. On doit alors définir deux types de rendements (que l'on appelle efficacités car ils sont supérieurs à 1) selon le type de machine. S'il s'agit d'un réfrigérateur, les grandeurs intéressantes sont le travail W que l'on dépense et le transfert thermique pris à la source froide (intérieur du réfrigérateur), soit $e = \frac{|Q_2|}{|W|} = \frac{Q_2}{W}$; tandis que s'il s'agit d'une pompe à chaleur, les grandeurs intéressantes sont le travail W et le transfert thermique Q_1 cédé à la source chaude (pièce à chauffer), soit $e = \frac{|Q_1|}{|W|} = \frac{-Q_1}{W}$. Dans les deux cas, le lecteur pourra s'exercer à exprimer les efficacités en fonctions des températures T_1 et T_2 grâce à la relation (37.11). On reviendra sur la notion de rendement dans le chapitre sur les machines thermiques.

3. Travail fourni par un opérateur lors du déplacement d'un piston

On s'intéresse au système décrit sur la figure 37.4.

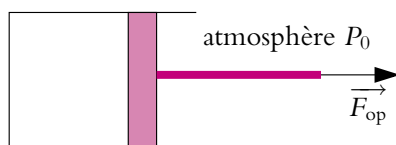


Figure 37.4 Travail de l'opérateur pendant le déplacement d'un piston.

Le système considéré est constitué du gaz dans le cylindre, du piston et du cylindre, mais comme cela a été décrit dans le paragraphe 3.1, on néglige l'énergie interne de l'enceinte et du piston, la variation d'énergie cinétique de ce dernier étant nulle. Lorsqu'un piston se déplace (par exemple dans un moteur ou un compresseur), il est soumis à deux forces extérieures : celle exercée par l'atmosphère et celle qui permet le déplacement du piston qu'on appelle *force de l'opérateur*. C'est le travail de cette dernière (travail dit *travail de l'opérateur*) qui sera coûteux dans le cas d'un compresseur et qui sera récupérable dans le cas d'un moteur. Ce travail n'est pas fourni par un opérateur à proprement parler mais par les parties de la machine autres que le cylindre et le piston considérés.

$$\text{État initial (1)} \left| \begin{array}{c} V_1 \\ P_1 \\ T_1 \end{array} \right| \quad \text{État final (2)} \left| \begin{array}{c} V_2 \\ P_2 \\ T_2 \end{array} \right|$$

On applique le premier principe au système (S) soit :

$$\Delta U_{\text{gaz}} + \Delta E_{\text{cylindre}} + \Delta U_{\text{piston}} + \Delta U_{\text{enceinte}} \simeq \Delta U_{\text{gaz}} = W + Q$$

Or le travail se décompose en deux termes : le travail de l'atmosphère et le travail de l'opérateur : $W = W_{\text{atm}} + W_{\text{op}} = P_0 V_{\text{balayé}} + W_{\text{op}} = P_0(V_1 - V_2) + W_{\text{op}}$ soit en combinant les deux relations :

$$W_{\text{op}} = \Delta U - Q - P_0(V_1 - V_2) \quad (37.14)$$

C'est normalement ce travail qui devrait intervenir dans un calcul de rendement puisque celui de l'atmosphère est gratuit et non récupérable.

Cependant il arrive souvent que l'on confonde W et W_{op} car le travail de l'atmosphère est en général négligeable devant celui de l'opérateur.

4. Détente de Joule-Gay Lussac

On appelle *détente* une transformation au cours de laquelle la pression diminue.

4.1 Description de l'expérience

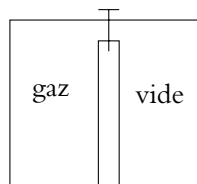


Figure 37.5 État initial de la détente de Joule - Gay Lussac.

Le système considéré est un gaz initialement dans un compartiment fermé et indéformable (compartiment gauche de la figure 37.5). Les parois des deux enceintes sont indéformables et adiabatiques. Le compartiment de droite est initialement vide. À l'instant $t = 0$, on ouvre le robinet permettant la communication entre les deux compartiments. Quelle est l'évolution du système ?

Système : gaz + les 2 enceintes

$$\text{État initial (1) Gaz} \left| \begin{array}{c} V_1 \\ P_1 \\ T_1 \end{array} \right| \quad \parallel \quad \text{État final (2) Gaz} \left| \begin{array}{c} V_2 \\ P_2 \\ T_2 \end{array} \right|$$

On calcule la variation d'énergie interne ΔU .

- Les parois sont adiabatiques donc le transfert thermique reçu par le système est nul : $Q = 0$.
- Les parois sont indéformables donc le travail reçu par le système de la part de l'extérieur est nul : $W = 0$.

La question que l'on peut se poser est : « le gaz qui déjà entré dans l'enceinte de droite va-t-il s'opposer à l'entrée du gaz encore dans l'enceinte de gauche ? ». C'est vrai mais, dans le calcul du travail, ce sont les forces extérieures qui interviennent et non les forces intérieures dont le travail est inclus dans l'énergie interne (voir paragraphe 1.1).

Ainsi si on applique le premier principe entre l'état initial et l'état final :

$$\Delta U = \Delta U_{\text{gaz}} + \Delta U_{\text{enceintes}} \simeq \Delta U_{\text{gaz}} = W + Q = 0 \quad (37.15)$$

La détente de Joule - Gay Lussac est une transformation à énergie interne constante.

D'autre part, il n'est pas besoin de s'étendre sur le fait que le gaz ne reviendra pas spontanément dans le compartiment de droite. **La transformation est irréversible.**

► **Remarque :** En ouvrant le robinet, on a déplacé une paroi mais le travail associé est négligeable.

4.2 Cas du gaz parfait

D'après la première loi de Joule (35.17), l'énergie interne d'un gaz parfait ne dépend que de la température. Puisqu'une détente de Joule - Gay Lussac est à énergie interne constante, elle est aussi à température constante :

$$\Delta U = 0 \Leftrightarrow C_V \Delta T = 0 \Leftrightarrow \Delta T = 0 \Leftrightarrow T_2 = T_1 \quad (37.16)$$

Pour un gaz parfait, une détente de Joule - Gay Lussac a lieu à température constante.

De plus, on peut dire qu'un gaz qui subit une détente de Joule - Gay Lussac sans variation de température, suit la première loi de Joule (35.17), c'est-à-dire que si $\Delta U = 0$ et $\Delta T = 0$, alors l'énergie interne U n'est fonction que de la température T .

4.3 Cas du gaz de Van der Waals

L'expression de l'énergie interne du gaz de Van der Waals a déjà été donnée (35.19), a étant une constante positive :

$$U = C_V T - \frac{n^2 a}{V}$$

Si on fait subir une détente de Joule - Gay Lussac à un gaz de Van der Waals, en prenant les mêmes notations pour les états initial et final qu'au paragraphe précédent, on a :

$$\Delta U = 0 \Leftrightarrow C_V(T_2 - T_1) = n^2 a \left(\frac{1}{V_2} - \frac{1}{V_1} \right) \quad (37.17)$$

or puisque $V_1 < V_2$, on en déduit que $T_2 < T_1$. Ainsi la détente de Joule - Gay Lussac s'accompagne d'une diminution de température.

On peut se convaincre de la véracité de ce résultat si on possède une bouteille pour fabriquer de la crème Chantilly à l'aide de cartouche d'air comprimé. Lorsqu'on perce cette cartouche, l'air comprimé initialement se détend dans la bouteille et la cartouche devient beaucoup plus froide qu'elle n'était initialement.

En pratique, la plupart des gaz étant bien décrits, aux températures usuelles, par l'équation de Van der Waals, on observe effectivement une diminution de température lors d'une détente de Joule - Gay Lussac.

Cependant, il existe en particulier une exception pour le dihydrogène H_2 aux fortes pressions, pour lequel on observe un réchauffement. En fait, l'équation de Van der Waals ne convient pas pour le dihydrogène aux fortes pressions.

La mesure de la variation de température permet de mesurer le coefficient a d'un gaz de Van der Waals grâce à l'équation (37.17)

Peut-on interpréter physiquement la diminution de température ?

On a vu que la différence entre un gaz parfait et un gaz de Van der Waals provenait de la prise en compte des forces d'interaction entre molécules appelées à juste titre « forces de Van der Waals ». L'énergie potentielle dont dérivent ces forces est de la forme $E_{pM} = -\frac{k}{r^n}$, où k est une constante positive et n un entier naturel, r représentant la distance moyenne entre molécule du gaz. Or l'énergie interne U est telle que $U = E_{pM} + E_{cM}$ (voir paragraphe 1.1)

Ainsi, dans cette transformation au cours de laquelle le volume augmente, la distance r moyenne entre particule augmente elle aussi, donc E_{pM} augmente aussi. Or l'énergie interne U est constante, donc E_{cM} diminue, et puisque l'énergie cinétique microscopique est proportionnelle à la température T , cette dernière diminue aussi.

On peut voir le problème de manière légèrement différente : en moyenne, les molécules sont plus éloignées les unes des autres dans l'état final que dans l'état initial. Il a fallu donc fournir du travail au gaz pour vaincre les interactions attractives entre molécules. Le système étant isolé, cette énergie ne peut provenir que de l'énergie cinétique des molécules qui diminue donc, de même que la température.

Dans le cas du dihydrogène H_2 sous forte pression, les interactions entre les molécules peuvent devenir répulsives car elles peuvent s'approcher très près les unes des autres à cause de leur petite taille ; les effets concernant l'énergie potentielle d'interaction sont inversés par rapport au cas précédent, ce qui se traduit par un gain d'énergie cinétique microscopique donc par une augmentation de température.

5. Détente de Joule - Thomson ou Joule - Kelvin

5.1 Étude de la détente

On appelle la transformation décrite dans ce paragraphe indifféremment détente de Joule - Thomson ou Joule - Kelvin car le physicien Thomson est devenu Lord Kelvin après avoir été anobli.

a) Dispositif

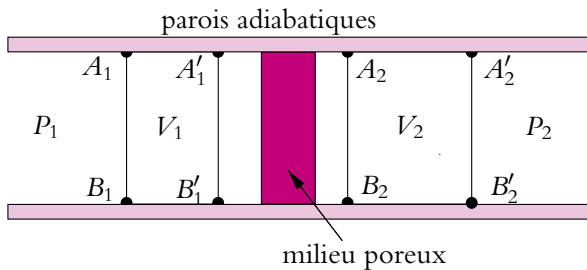


Figure 37.6 Dispositif utilisé pour la détente de Joule - Thomson.

On étudie l'écoulement lent d'un fluide dans un tube séparé en deux régions par un milieu poreux (coton, verre fritté...).

- Le tube est formé de **parois indéformables et adiabatiques**.
- Dans la partie gauche, la pression est uniforme et égale à P_1 et dans la partie droite, la pression est aussi uniforme mais égale à P_2 . Pour qu'il y ait écoulement de gauche à droite, il faut que $P_1 > P_2$.

- On suppose que **l'écoulement est stationnaire**, c'est-à-dire que toutes les variables macroscopiques sont indépendantes du temps.
- On suppose de plus que l'écoulement est lent ce qui permet de négliger l'énergie cinétique macroscopique du fluide.

b) Définition des systèmes

On a affaire à un écoulement : **c'est un système ouvert**. Il faut **se ramener à un système fermé** pour pouvoir appliquer le premier principe.

Étant donné que l'écoulement est permanent, l'état d'une tranche de gaz ne dépend que de sa position : c'est-à-dire que deux tranches de gaz différentes occupant la même position mais à deux instants différents t_1 et t_2 seront dans le même état thermodynamique lorsqu'elles occupent cette position.

On considère le système fermé (S') à cheval sur le milieu poreux : tranche $A_1A_2B_1B_2$ (notée A_1A_2 dans la suite) qui après écoulement (pendant dt) devient $A'_1A'_2B'_1B'_2$ (notée $A'_1A'_2$ dans la suite).

On aura besoin d'un autre système (S) : $A_1A'_1B_1B'_1$ (noté dans la suite $A_1A'_1$) qui devient après passage à travers la zone poreuse $A_2A'_2B_2B'_2$ (noté dans la suite $A_2A'_2$). Pour le système (S), on pose :

État initial (1)	Volume V_1 Pression P_1 Énergie interne U_1		État final (2)	Volume V_2 Pression P_2 Énergie interne U_2
------------------	--	--	----------------	--

c) Premier principe

On applique le premier principe à (S') entre l'état initial et l'état final :

$$\Delta U_{(S')} = W + Q$$

avec $Q = 0$ puisque la transformation est adiabatique.

Comment calculer le travail W ? Les parois sont indéformables donc seules interviennent les forces de pression du reste du gaz.

À gauche, la limite A_1B_1 s'est déplacée d'une longueur L_1 et a balayé un volume V_1 . De ce côté, le système est **poussé** par une force $F_1 = P_1\Sigma$ (voir figure 37.6) où Σ est la section. Le travail reçu est donc **moteur** et vaut :

$$W_1 = F_1L_1 = P_1\Sigma L_1 = P_1V_1$$

Comme on l'a signalé au paragraphe 3.1, on retrouve ici le volume balayé qui est différent de la variation de volume. À droite, le raisonnement est identique : la limite

A_2B_2 s'est déplacé d'une longueur L_2 et a balayé un volume V_2 . De ce côté, le système est **retenu** par une force $F_2 = P_2 \Sigma$. Le travail reçu est donc **résistant** et vaut :

$$W_2 = -F_2 L_2 = -P_2 \Sigma L_2 = -P_2 V_2$$

On en déduit le travail total reçu par le système (S') :

$$W = W_1 + W_2 = P_1 V_1 - P_2 V_2 \quad (37.18)$$

d) Enthalpie constante

L'expression du travail fait intervenir les variables du système (S). On est donc amené à effectuer un changement de système pour montrer que les variations d'énergie interne de (S') et (S) sont égales. On compare alors les systèmes (S) et (S') en faisant intervenir la partie commune centrale $A'_1 A_2 B'_1 B_2$ sur le dessin. On a déjà signalé que comme l'écoulement est indépendant du temps, l'énergie de cette partie reste constante. Ainsi :

- $A_1 A_2 B_1 B_2 = A_1 A'_1 B_1 B'_1 + A'_1 A_2 B'_1 B_2$
- $A'_1 A'_2 B'_1 B'_2 = A'_1 A_2 B'_1 B_2 + A_2 A'_2 B_2 B'_2$

Si on calcule la variation d'énergie interne de (S) et (S') en faisant intervenir la partie centrale, on trouve :

$$\Delta U_{(S')} = U_{A'_1 A'_2 B'_1 B'_2} - U_{A_1 A_2 B_1 B_2} = (U_2 + U_{A'_1 A_2 B'_1 B_2}) - (U_1 + U_{A'_1 A_2 B'_1 B_2})$$

et après simplification :

$$\Delta U_{(S')} = U_2 - U_1 = \Delta U_{(S)} \quad (37.19)$$

Les deux systèmes (S) et (S') ont même variation d'énergie interne. Le système (S') est plus aisé à utiliser pour calculer le travail reçu mais ce sont les variables de (S) qui apparaissent naturellement.

En appliquant le premier principe à (S'), on obtient d'après (37.19) :

$$\Delta U_{(S')} = \Delta U_{(S)} = U_2 - U_1 = W + Q$$

soit

$$U_2 - U_1 = P_1 V_1 - P_2 V_2 \quad (37.20)$$

en regroupant les termes :

$$(U_2 + P_2 V_2) - (U_1 + P_1 V_1) = 0 \Leftrightarrow \Delta H = 0$$

La détente de Joule - Thomson est une transformation à enthalpie constante : on dit qu'elle est *isenthalpique*.

D'autre part, le fluide ne peut revenir seul en arrière, puisqu'il se déplace spontanément dans le sens des pressions décroissantes. La transformation est donc **irréversible**

► **Remarque :** On étudiera dans le chapitre sur les machines thermiques le cas où l'écoulement n'est pas horizontal et de vitesse non négligeable.

5.2 Cas d'un gaz parfait

Un gaz parfait obéit à la deuxième loi de Joule (36.20) donc si sa variation d'enthalpie est nulle, sa variation de température l'est aussi.

$$\Delta H = 0 \Leftrightarrow C_p \Delta T = 0 \Leftrightarrow \Delta T = 0 \Leftrightarrow T_2 = T_1 \quad (37.21)$$

Un gaz parfait subit une détente de Joule - Thomson isotherme, ce qui est équivalent à la deuxième loi de Joule. En effet, si $\Delta H = 0$ et $\Delta T = 0$, l'enthalpie H n'est fonction que de la température T .

5.3 Cas d'un gaz de Van der Waals

Pour un gaz de Van der Waals, l'évolution est plus difficile à prévoir ; elle dépend de la température initiale du gaz. Si on étudie le gaz dans un diagramme d'Amagat, on s'aperçoit qu'il existe deux zones séparées par une parabole (représentée en trait pointillé (Cf. figure 37.7)).

Si le gaz se trouve initialement dans la zone à l'intérieur de la parabole, il se refroidit, sinon il se réchauffe. On peut tracer, sur le diagramme, une isotherme T_i tangente à la parabole. On appelle cette température T_i , la *température d'inversion*. Il est donc nécessaire que la température du fluide soit inférieure à T_i pour pouvoir le refroidir par une détente de Joule - Thomson.

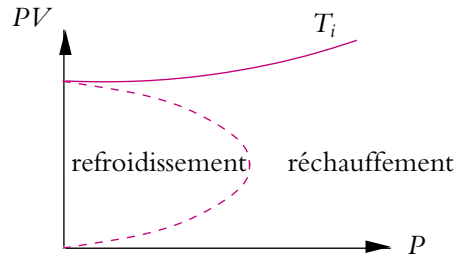


Figure 37.7 Diagramme d'Amagat.

Corps	T_i
He	34 K
H ₂	202 K
N ₂	625 K

Tableau 37.1 Températures d'inversion de différents corps.

Comme le montre la table 37.1, le diazote peut être refroidi par une détente de Joule - Thomson en partant d'une température ambiante. Par contre, pour l'hélium et le dihydrogène, il sera nécessaire de les refroidir auparavant par une autre méthode : détente de Joule - Gay Lussac ou mise en contact avec un corps plus froid. Ces méthodes peuvent être utilisées pour refroidir les gaz afin de les liquéfier.

6. Récapitulatif : utiliser U ou H

On récapitule dans quels cas il faut utiliser l'énergie interne et dans quels cas il faut utiliser l'enthalpie pour appliquer le premier principe.

GAZ	
U	H
Isochore : $\Delta U = Q + W_{\text{autre}}$ Exemples : détente Joule - Gay-Lussac. Fluide mis dans un récipient vide indéformable	Isobare ou monobare entre deux états d'équilibre : $\Delta H = Q + W_{\text{autre}}$ Exemples : détente de Joule - Thomson. Écoulement où le travail des forces de pression est comptabilisé dans ΔH

Dans les autres cas, on essaiera U dans un premier temps et on passera à H si nécessaire. Pour les solides et les liquides peu compressibles, on a vu que les variations de U et H sont pratiquement confondues on peut donc utiliser l'une ou l'autre.

A. Applications directes du cours

1. Vrai ou faux ?

Justifier les réponses et rectifier les affirmations fausses si c'est possible.

1. L'énergie interne d'un gaz ne dépend que de la température.
2. Au cours d'une transformation monobare, le transfert thermique est égal à la variation d'enthalpie.
3. La température d'un gaz réel diminue au cours d'une détente de Joule-Gay-Lussac.
4. La température d'un gaz réel diminue au cours d'une détente de Joule-Thomson.
5. La température d'un corps pur augmente quand on lui fournit un transfert thermique.
6. La température d'un système isolé reste constante.
7. Un système fournit adiabatiquement du travail.
 - a) Sa température ne varie pas.
 - b) Son énergie interne diminue.
8. Un gaz est contenu dans un cylindre fermé par un piston. On diminue son volume de moitié de manière soit adiabatique, soit isotherme. Le travail à fournir est :
 - a) plus grand si la compression est adiabatique que si elle est isotherme ;
 - b) égal à la variation d'énergie interne dans les deux cas.

2. Valeur en eau d'un calorimètre

On mélange 95 g d'eau à 20°C et 71 g d'eau à 50°C dans un calorimètre.

1. Quelle est la température du mélange final à l'équilibre en négligeant l'influence du calorimètre ?
2. Expérimentalement on obtient 31,3°C. Expliquer.
3. En déduire la valeur en eau du calorimètre.

3. Détermination de la chaleur massique d'un métal

Dans un calorimètre dont la valeur en eau est de 41 g, on verse 100 g d'eau. Une fois l'équilibre thermique atteint, on mesure une température de 20°C. On plonge alors une barre métallique dont la masse est 200 g et dont la température initiale est de 60°C. A l'équilibre, on mesure une température de 30°C. Déterminer la capacité thermique massique du métal.

On donne la capacité thermique massique de l'eau $c_e = 4,18 \text{ kJ} \cdot \text{kg}^{-1} \cdot \text{K}^{-1}$ et on suppose que les capacités thermiques massiques sont constantes dans le domaine de températures considérées.

4. Transformations polytropiques

Une transformation polytropique est une transformation quasistatique vérifiant PV^k constante.

1. Calculer le travail des forces de pression pour un gaz parfait subissant une transformation polytropique entre (P_0, V_0, T_0) et (P_1, V_1, T_1) en fonction des pressions et volumes ainsi que de k .

2. On note $\gamma = \frac{C_p}{C_v}$ qui est une constante pour un gaz parfait. Déterminer l'expression du transfert thermique au cours de la transformation précédente sous la forme $C(T_1 - T_0)$.
3. Donner une interprétation physique de C .
4. Étudier en les interprétant physiquement chacun des cas suivants :
 - a) $k = \gamma$,
 - b) $k = 0$,
 - c) $k \rightarrow \infty$,
 - d) $k = 1$.
5. Pour les quatre cas précédents, peut-on retrouver facilement les résultats classiques des travaux correspondants ?

5. Transformations d'un gaz de Van der Waals

Un gaz réel suit l'équation de Van der Waals :

$$\left(P + \frac{a}{V^2}\right)(V - b) = RT$$

Il subit la succession de transformations suivantes qui seront supposées quasistatiques :

A B dilatation isobare à la pression P_0 augmentant le volume de $V_0 = 1 \text{ L}$ à $V_1 = 5 \text{ L}$,

B C compression isotherme à la température de 1000 K ramenant le volume à V_0 ,

C A détente isochore ramenant le système à son état initial.

1. Tracer ce cycle dans le diagramme de Watt.
2. Dans les unités du système international, $a = 0,14$ et $b = 3,22 \cdot 10^{-5}$. Préciser ces unités.
3. Préciser les valeurs de la température, de la pression et du volume aux sommets du cycle.
4. Calculer les travaux des forces de pression au cours de chaque phase du cycle.
5. Au cours du cycle, le système reçoit-il ou cède-t-il un transfert thermique ? Justifier numériquement.
6. Déterminer les transferts thermiques au cours de chaque phase du cycle.
7. Quelle vérification peut-on faire ?

On donne l'expression de son énergie interne

$$U = \frac{3}{2}RT - \frac{a}{V}$$

6. Compressions et détente successives d'un gaz parfait

On étudie les compressions et détente successives d'un gaz parfait ($\gamma = 1,4$). On suppose que les transformations sont quasistatiques.

1. Le gaz subit une compression isotherme (température $T_0 = 273 \text{ K}$) de $P_0 = 1 \text{ bar}$ à $P_1 = 20 \text{ bar}$, puis une détente adiabatique jusqu'à P_0 . Calculer la température T_1 après les deux transformations.

- On recommence l'opération. Calculer T_2 puis T_n après n séries de transformations. Effectuer l'application numérique pour $n = 5$.
- Calculer la variation d'énergie interne au cours de la $n^{\text{ème}}$ double transformation en fonction de T_0 , P_0 et P_1 . Calculer le travail W et le transfert thermique Q reçus par le gaz. Effectuer les applications numériques pour une mole de gaz et pour $n = 5$.

B. Exercices et problèmes

1. Chauffage d'une enceinte

On étudie le système suivant :

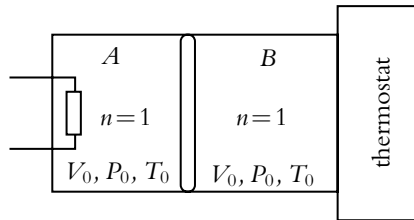


Figure 37.8

On suppose que les enceintes contiennent des gaz parfaits et que l'enceinte A est parfaitement calorifugée. On note $\gamma = \frac{C_p}{C_v}$. On chauffe l'enceinte A jusqu'à la température T_1 par la résistance chauffante. Les transformations seront considérées comme quasistatiques.

- Déterminer les volumes finaux des deux enceintes ainsi que la pression finale.
- Calculer la variation d'énergie interne de chacune des deux enceintes ainsi que celle de l'ensemble.
- Quelle est la nature de la transformation de l'enceinte B ? En déduire le calcul du travail échangé entre les deux enceintes et le transfert thermique Q_1 entre l'enceinte B et le thermostat.
- Déterminer le transfert thermique Q_2 fourni par la résistance.

2. Compression quasistatique ou non

Un cylindre vertical de hauteur $h = 10$ cm à parois adiabatiques contient un gaz parfait dont on supposera le rapport $\gamma = 1,4$ indépendant de la température. Le cylindre est placé dans le vide, la pression du gaz étant équilibrée par un piston de masse m_0 , de section s , mobile sans frottement. La température initiale du gaz est $T_0 = 273$ K.

- On ajoute progressivement de petites masses (infiniment petites) sur le piston dont la somme $m = 2m_0$. Exprimer, en fonction des données, la hauteur h_1 dont est descendu le piston une fois l'équilibre final réalisé, ainsi que la température finale T_1 .
- Partant du même état initial, on pose une masse m sur le piston, puis on la lâche. Exprimer, en fonction des données, la hauteur h_2 dont est descendu le piston une fois l'équilibre final réalisé ainsi que la température finale T_2 .

3. Étude d'un compresseur, d'après ESIM 2000

Le problème étudie le compresseur d'un moteur à air comprimé (celui d'un marteau-piqueur, par exemple). L'air est assimilé à un gaz parfait de masse molaire $M = 29 \text{ g.mol}^{-1}$, de capacité thermique massique à pression constante $c_p = 1,00 \text{ kJ.kg}^{-1}.\text{K}^{-1}$ et de rapport des capacités thermiques à pression et à volume constants $\gamma = 1,4$. La constante des gaz parfaits est $R = 8,314 \text{ J.K}^{-1}.\text{mol}^{-1}$.

L'air est aspiré dans les conditions atmosphériques, sous la pression $P_0 = 1 \text{ bar}$ et à la température $T_0 = 290 \text{ K}$, jusqu'au volume V_m , puis comprimé jusqu'à la pression P_1 , où il occupe le volume V_1 , et refoulé à la température T_1 dans un milieu où la pression est $P_1 = 6 \text{ bar}$. Bien que le mécanisme réel d'un compresseur soit différent, on suppose que celui-ci fonctionne comme une pompe à piston, qui se compose d'un cylindre, d'un piston coulissant entraîné par un moteur et de deux soupapes.

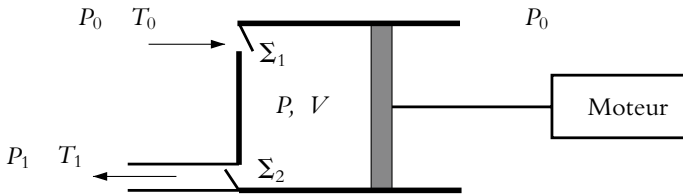


Figure 37.9

- La soupape d'entrée Σ_1 est ouverte si la pression P dans le corps de pompe est inférieure ou égale à la pression atmosphérique P_0 .
- La soupape de sortie Σ_2 est ouverte si P est supérieure à P_1 .
- Le volume V du corps de pompe est compris entre 0 et V_m .
- À chaque cycle (chaque aller et retour du piston), la pompe aspire et refoule une mole d'air.

1. a) Tracer sur un diagramme de Watt (P en ordonnée, V en abscisse) l'allure de la courbe représentant un aller et un retour du piston. Indiquer le sens de parcours par une flèche.

b) Montrer que le travail de l'air à droite du piston est nul sur un aller-retour. Représenter sur ce diagramme le travail fourni par le moteur qui actionne le piston. On supposera que le mouvement est assez lent pour que l'évolution soit réversible.

2. Pendant la phase de compression, l'air suit une loi polytropique $PV^k = \text{cste}$; il sort du compresseur à la température $T_1 = 391 \text{ K}$.

a) Montrer que $k = 1,20$.

b) Exprimer le travail mécanique W_{moteur} fourni par le moteur pendant un aller-retour en fonction de R , n , k , T_1 et T_0 . Calculer numériquement W_{moteur} .

c) Au cours d'un cycle, le système constitué par une mole d'air est transvasé et reçoit le transfert thermique Q . En tenant compte du travail des forces de pression du gaz extérieur au cylindre de la pompe, démontrer que la somme $W_{\text{moteur}} + Q = \Delta H$. Calculer numériquement Q . Interpréter le signe de Q .

3. Le débit est de $0,013 \text{ kg.s}^{-1}$. Calculer la puissance P_{moteur} fournie par le moteur.

4. Moteur de scooter, d'après ENSTIM 1999

Le moteur d'un scooter est un moteur à deux temps :

- Premier temps : compression du mélange air-carburant (FB), puis combustion (BC).
- Deuxième temps : détente (CD), échappement des gaz de combustion et admission d'une nouvelle charge de mélange carburé (DHF).

On fait les hypothèses suivantes :

- la combustion (BC) est instantanée et le piston n'a pas le temps de se déplacer : $V = V_C$;
- la détente (CD) et la compression (FB) sont adiabatiques réversibles ;
- lors de l'échappement et de l'admission (DF) quasi instantanés, le volume du cylindre est constant : $V = V_D$.

On appelle cylindrée du moteur la différence $V_D - V_C$ et taux de compression le rapport $a = \frac{V_D}{V_C}$. On suppose que le mélange air-carburant est un gaz parfait ($\gamma = 1,4$) de masse molaire $M = 29 \text{ g.mol}^{-1}$. La constante des gaz parfaits est : $R = 8,31 \text{ J.mol}^{-1}.\text{K}^{-1}$.

En F : $T_F = 300 \text{ K}$ et $P_F = 10^5 \text{ Pa}$.

La notice technique d'un scooter donne :

- vitesse maximale : 45 km.h^{-1} ;
- régime de puissance maximale : $7000 \text{ tours.min}^{-1}$;
- puissance maximale : $4,40 \text{ kW}$;
- cylindrée : $49,5 \text{ cm}^3$;
- course du piston : $39,2 \text{ mm}$.

1. Tracer l'allure du cycle ($FBCDF$) sur un diagramme de Watt (P, V).

2. Lorsque le scooter roule à son allure maximale (le régime est aussi maximal), quelle est la durée d'un cycle ?

Comparer cette vitesse à un ordre de grandeur de la vitesse quadratique moyenne des molécules en F .

Peut-on en déduire une caractéristique des transformations (CD) et (FB) ?

3. La pression en fin de compression est de $6 \cdot 10^5 \text{ Pa}$. En déduire le taux de compression a .

4. Exprimer le travail fourni W par le moteur au cours d'un cycle en fonction de R, γ, P_F, V_D et des températures T_F, T_B, T_C et T_D .

5. Exprimer la chaleur Q_{BC} reçue par le système lors de la combustion, en fonction de $R, \gamma, P_F, V_D, T_C, T_B, T_F$.

6. On définit le rendement du moteur par : $\eta = \left| \frac{W}{Q_{BC}} \right|$. Exprimer η en fonction de a et γ .

7. En prenant $\eta = 0,4$, calculer la chaleur libérée à chaque cycle par la combustion lorsque le scooter roule à 45 km.h^{-1} , sa puissance étant $P = 4,4 \text{ kW}$.

8. Sachant que le pouvoir calorifique de l'essence, c'est-à-dire la chaleur libérée par la combustion, est $q = 30 \text{ kJ.cm}^{-3}$, déterminer la consommation pour 100 km .

5. Mesure du coefficient γ , d'après INA-ENSA 2000

1. Méthode de Rüchardt :

La méthode de Rüchardt permet de déterminer le rapport $\gamma = \frac{C_P}{C_V}$ des capacités thermiques à pression et à volume constants en étudiant le mouvement d'une bille de masse $m = 16,6$ g dans un tube en verre. La bille métallique dont le diamètre est très voisin de celui du tube se comporte comme un piston étanche. On place le tube au-dessus d'un récipient de volume $V_0 = 10$ L.

Lorsqu'on lâche la bille dans le tube de section $s = 2,0 \cdot 10^{-4}$ m², on observe des oscillations amorties autour d'une position d'équilibre. La méthode consiste à mesurer la période d'oscillation T_0 du mouvement de la bille dans le tube. Pour cela, on enregistre la pression P dans le récipient à l'aide d'un capteur de pression pendant 20 secondes. On observe qu'à l'équilibre, la position de la bille est $z_e = -41$ cm. On prendra pour axe Oz un axe vertical ascendant dont l'origine est à l'extrémité supérieure du tube.

On note $P_0 = 1,0 \cdot 10^5$ Pa la pression atmosphérique, $g = 9,81$ m.s⁻² l'accélération de la pesanteur, C_P la capacité thermique molaire à pression constante, C_V la capacité thermique molaire à volume constant, $T_0 = 293$ K la température extérieure, T la température à un instant donné du gaz contenu dans le récipient, z la position de la bille à un instant donné, $R = 8,31$ J.K⁻¹.mol⁻¹ la constante des gaz parfaits, $M = 29$ g.mol⁻¹ la masse molaire de l'air et $1,0$ bar = $1,0 \cdot 10^5$ Pa.

- Déterminer l'ordre de grandeur de la suppression créée par la bille. Le comparer à la valeur de la pression atmosphérique. Que peut-on en conclure ?
- Pour mesurer la pression, on utilise un capteur de pression commercial qui délivre une tension u proportionnelle à la pression P soit $u = \alpha P$. En l'absence de bille, la pression est mesurée à l'aide d'un multimètre numérique. Sur le calibre le plus approprié, on obtient 3,333 V. Quelle est la valeur de α en V.bar⁻¹ ? En déduire l'ordre de grandeur des variations de la tension u . Conclure quant à l'utilisation du multimètre pour cette mesure.
- L'enregistrement de l'écart de pression $P - P_0$ a été effectué à l'aide du montage ci-dessous.

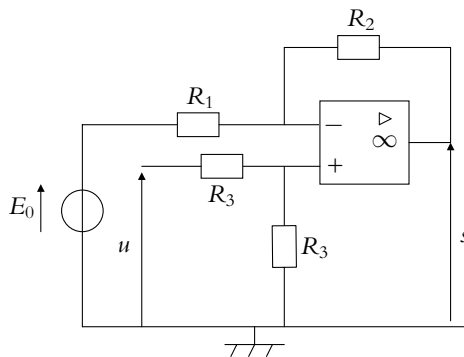


Figure 37.10

On suppose que l'amplificateur opérationnel est idéal et fonctionne en régime linéaire. Exprimer la tension de sortie s de l'amplificateur opérationnel en fonction de u , R_1 , R_2 et E_0 .

- d) Montrer qu'il s'agit d'une fonction affine $s = aP + b$ de la pression P exprimée en bars. On précisera les expressions de a et b .
- e) Sachant que la tension de sortie s de l'amplificateur opérationnel est nulle lorsque la pression P vaut la pression atmosphérique et qu'on a $s = 3,0$ V pour un écart de pression $P - P_0$ de 1,0 kPa, déterminer les valeurs de R_2 et E_0 . On donne $R_1 = 1,0$ k Ω .
- f) Montrer que la force de pression exercée sur une demi-sphère de rayon r par une pression uniforme P_0 est $F = P_0 \pi r^2$. On donnera la direction de cette force.
- g) En appliquant par exemple le principe fondamental de la dynamique à la bille, établir l'équation du mouvement de la bille en fonction de P , P_0 , m , g , s et des dérivées temporelles de z .
- h) L'air contenu dans le récipient est assimilé à un gaz parfait. On suppose qu'il vérifie la relation de Laplace disant que le produit PV^γ est constant. Rappeler les conditions de validité de cette loi.
- i) En admettant que $dV = V - V_0 = sz$ et que $dP = P - P_0$, en déduire $P - P_0$ en fonction de z .
- j) En déduire l'équation différentielle du mouvement vertical de la bille et montrer que la pression obéit à une équation différentielle analogue.
- k) En déduire l'expression de la période T_0 des oscillations et celle de la position z_e à l'équilibre.
- l) En déduire l'expression de γ en fonction de T_0 .
- m) Sachant que $T_0 = 1,12$ s, déterminer numériquement γ .
- n) La position d'équilibre observée est-elle compatible avec l'hypothèse précédente ?

2. Méthode de Clément-Desormes :

Pour un fluide quelconque, on peut écrire le transfert thermique sous la forme :

$$\delta Q = C_V dT + l dV = C_P dT + h dP \quad \text{avec} \quad l = T \left(\frac{\partial P}{\partial T} \right)_V \quad \text{et} \quad h = -T \left(\frac{\partial V}{\partial T} \right)_P$$

- a) Définir les trois coefficients thermoélastiques α , β et χ_T .
- b) Établir l'expression de l et h en fonction des coefficients thermoélastiques ainsi que de T , P et V .
- c) En déduire $C_P - C_V$ en fonction des coefficients thermoélastiques ainsi que de T , P et V .
- d) Exprimer α , β , χ_T , l , h et $C_P - C_V$ pour un gaz parfait.
- e) Montrer que, dans un diagramme de Clapeyron donnant la pression P en fonction de V , la pente d'une adiabatique est γ fois plus élevée que celle d'une isotherme.
- f) On considère un récipient contenant de l'air et possédant deux orifices. L'un est relié au capteur de pression déjà utilisé dans la méthode de Rüchardt et l'autre à un robinet. Initialement on crée une légère surpression dans le récipient. En A , on ouvre le robinet qu'on ferme très rapidement en B tout en enregistrant la pression. On laisse ensuite le système évoluer

jusqu'en C . On observe une diminution de la tension de $u_A = 3,5 \text{ V}$ à $u_B = 0,1 \text{ V}$ puis une augmentation jusqu'à $u_C = 0,7 \text{ V}$.

Quel système doit-on considérer ?

g) Pourquoi peut-on supposer que la transformation AB est adiabatique ? On fera également l'hypothèse qu'elle est quasistatique, ce qu'on ne demande pas de justifier.

h) Préciser la nature de la transformation BC .

i) Pour cette transformation BC , on admet que les échanges thermiques se font par convection selon la loi dite de Newton : les pertes sont proportionnelles à l'écart de température entre l'intérieur et l'extérieur $\mathcal{P} = k(T - T_0)$ où T est la température du gaz et T_0 la température extérieure. En faisant un bilan énergétique, montrer que $\frac{dT}{dt} = -\frac{T - T_0}{\tau}$. On exprimera τ en fonction de R , n le nombre de moles du gaz dans le système et γ .

j) En déduire l'équation différentielle vérifiée par la pression.

k) Tracer dans le diagramme de Clapeyron l'évolution du système.

l) Exprimer γ en fonction des tensions u_A , u_B et u_C . Donner sa valeur numérique.

Deuxième principe de la thermodynamique

38

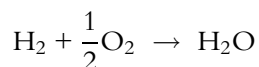
1. Nécessité d'un second principe. Limites du premier principe

On a introduit avec l'énoncé du premier principe les notions de transfert thermique et de travail mais rien ne différencie ces deux notions *a priori* qui ont d'ailleurs la même unité.

Lorsqu'on écrit pour un cycle $\Delta U = W + Q = 0$ soit $Q = -W$, il semble qu'il y ait équivalence entre Q et W . En particulier, le premier principe peut laisser penser que le moteur ($W < 0$) avec une seule source de chaleur ($Q > 0$) peut exister et qu'il a un rendement égal à 1 puisque $|W| = |Q|$. L'expérience prouve que ce moteur monotherme n'existe pas, ce que permet d'expliquer le deuxième principe (Cf. chapitre sur les machines thermiques).

D'autre part, rien n'interdit dans le premier principe d'inverser le sens d'une machine thermique, c'est-à-dire de pouvoir lui faire décrire un cycle dans le sens horaire ou trigonométrique : autrement dit, d'après le premier principe, il serait possible d'utiliser un réfrigérateur pour faire avancer une voiture.

On s'intéresse maintenant à une réaction chimique comme :



Si on met en présence le dihydrogène et le dioxygène en proportions stœchiométriques et qu'on amorce la réaction, le premier principe appliqué à la chimie dit que l'on récupère un transfert thermique (égal à la variation d'enthalpie du système si la réaction est monobare) de 242 kJ.mol^{-1} . On peut alors penser qu'en fournissant la même quantité d'énergie à une mole d'eau, on peut récupérer du dihydrogène et du dioxygène spontanément. On sait bien que c'est impossible, et qu'il faut avoir recours à des procédés plus complexes. Encore une fois, le premier principe ne donne aucune information sur le caractère irréversible de cette réaction.

On peut donner encore trois exemples de phénomènes irréversibles que le premier principe ne permet pas d'interpréter. L'effet Joule est toujours un effet dissipatif et si on change le sens du courant, l'échange d'énergie entre le circuit et l'extérieur est toujours dans le même sens. Quant à la diffusion de particules (colorant dans un solvant, tache d'encre sur un buvard), elle a toujours lieu des zones les plus concentrées vers les zones les moins concentrées. Enfin, pour conclure ce paragraphe, les transferts thermiques se font spontanément des corps chauds vers les corps froids.

Le premier principe est insuffisant pour expliquer ces phénomènes et étudier la notion d'irréversibilité.

2. Réversibilité - Irréversibilité

Les transformations réversibles ont déjà été évoquées dans le chapitre sur le premier principe mais c'est grâce au deuxième principe que cette notion peut être développée. On donne ici une définition plus précise d'une transformation réversible :

Une transformation est dite *réversible* si une modification infiniment petite des paramètres extérieurs permet d'en inverser le sens.

Une transformation réelle est en général non réversible mais on peut parfois réaliser des transformations réelles dont les états successifs sont très voisins d'une transformation réversible qui apparaît comme **une transformation idéale limite**. Mais pourquoi s'intéresse-t-on à des transformations réversibles ? Tout d'abord, pour une machine fonctionnant avec de telles transformations qui sont idéales, les pertes sont minimisées et le rendement augmenté. D'autre part, on a vu dans le chapitre sur le premier principe qu'il est souvent plus aisé de calculer les échanges d'énergie lorsque la transformation est réversible.

Ainsi, il arrive bien souvent que l'on fasse l'hypothèse qu'une machine est réversible sachant que le rendement calculé est maximal (Cf. chapitre sur les machines thermiques), le rendement réel étant inférieur au rendement calculé.

Les principales causes d'irréversibilité sont :

- les frottements solides,
- les frottements fluides,
- les gradients de concentration ou de température qui donnent lieu aux phénomènes de diffusion,
- les réactions chimiques.

► **Remarque :** En ce qui concerne l'effet Joule, on peut modéliser l'effet par un frottement visqueux ressenti par les charges mobiles, mais ce n'est qu'un modèle.

Pourquoi les frottements rendent-ils les transformations irréversibles ? Une force de frottement est toujours résistive (Cf. chapitre sur l'énergie en mécanique) quel que soit le sens du mouvement, donc le frottement est indépendant du sens de la transformation (par exemple dans le déplacement d'un piston).

En ce qui concerne les frottements solides, il n'existe pas de transformations réversibles limites. En effet, il n'est pas possible d'annuler le frottement qui est décrit par les lois dites de Coulomb vues en seconde année ou en Sciences Industrielles. Quand un objet est en mouvement, le frottement est quasiment indépendant de la valeur de la vitesse et puisqu'un objet peut rester en équilibre sur un plan incliné, c'est qu'il existe un frottement qui l'empêche de glisser. Ainsi, même en déplaçant l'objet lentement, on ne peut annuler le frottement. La seule solution pour le diminuer est de lubrifier les surfaces en contact.

Les transformations avec frottements fluides possèdent une transformation réversible limite. On a vu en mécanique que la force de frottement fluide est de la forme $h\nu^a$, h étant une constante et a un coefficient dépendant de la vitesse (Voir cours de mécanique). Ainsi, si on effectue une transformation à très faible vitesse, le frottement est quasiment inexistant. Il existe d'ailleurs une expérience qui permet de vérifier ceci : on dépose une goutte d'encre à la surface d'un liquide très visqueux comme de la glycérine. On étale très lentement la tache d'encre avec une aiguille. En repassant ensuite l'aiguille très lentement par le même chemin en arrière, on arrive à reformer la tache d'encre initiale. La transformation est réversible.

Les phénomènes de diffusion sont des phénomènes irréversibles comme cela a été évoqué au paragraphe précédent. Il n'y a pas de transformations réversibles limites.

Enfin, certaines réactions chimiques sont réversibles, on utilise le terme de *réactions équilibrées* en chimie. Dans ce cas, elles sont dans un état voisin de l'équilibre et en modifiant légèrement des paramètres (concentrations, pression...), on déplace l'équilibre.

2.1 Énoncé du deuxième principe

Pour un système thermodynamique donné, il existe une fonction d'état S appelée *entropie* (du grec *ενθροπια* entropia signifiant qui revient en arrière). L'entropie possède les propriétés suivantes :

- S est extensive ;
- lors d'une évolution **adiabatique** d'un système fermé : $\Delta S \geq 0$;
- lors d'une évolution **quelconque** d'un système fermé, on peut décomposer ΔS en $\Delta S = S_{\text{éch}} + S_{\text{créée}}$. Dans cette décomposition, $S_{\text{éch}}$ provient des transferts

thermiques et $S_{\text{créée}}$ est un terme de *création d'entropie* qui est toujours positif ou nul.

- $S_{\text{créée}} = 0$ pour une transformation réversible,
- $S_{\text{créée}} > 0$ pour une transformation irréversible.

L'expression générale de $S_{\text{éch}}$ est :

$$S_{\text{éch}} = \int \frac{\delta Q}{T_{\Sigma}} \quad (38.1)$$

où T_{Σ} est la température de la surface du système.

Il apparaît donc que contrairement au premier principe qui est un principe de **conservation**, le deuxième principe est un principe **d'évolution**.

L'entropie s'exprime en J.K^{-1} .

2.2 Conséquences immédiates

Le travail n'a pas d'effet direct sur l'entropie. Contrairement au premier principe, le deuxième principe différencie W et Q .

Entre deux états voisins, la variation élémentaire d'entropie s'écrira :

$$dS = \delta S_{\text{créée}} + \delta S_{\text{éch}}$$

Il faut noter qu'on écrit $\delta S_{\text{créée}}$ et $\delta S_{\text{éch}}$ et non $dS_{\text{créée}}$ et $dS_{\text{éch}}$, comme pour δQ et δW car ce ne sont pas des variations d'une fonction d'état entre deux états infiniment proches mais des petites quantités échangées ou créées.

Pour une évolution non adiabatique, ΔS peut être négative, si $S_{\text{éch}} < 0$ et $|S_{\text{éch}}| > S_{\text{créée}}$.

Dans le cas d'une transformation adiabatique réversible, si l'évolution de l'état noté (1) à l'état noté (2) est possible alors l'évolution inverse l'est aussi. Dans ce cas, on peut écrire :

$$\text{Évolution } 2 \rightarrow 1 : \Delta S = S_2 - S_1 \geq 0$$

$$\text{Évolution } 1 \rightarrow 2 : \Delta S' = S_1 - S_2 \geq 0$$

d'où :

$$\Delta S = 0$$

Une évolution adiabatique réversible est isentropique et réciproquement.

Lors d'une évolution adiabatique non réversible, $\Delta S > 0$ donc S croît au cours de l'évolution. Si la transformation s'arrête, c'est que l'entropie ne peut plus croître et

donc **S est maximale à l'équilibre**. Ainsi, l'Univers étant un système isolé, son entropie augmente au cours du temps.

3. Température et pression thermodynamique

Plusieurs définitions de la température et de la pression ont déjà données. En voici une dernière qui est bien sûr cohérente avec les autres, ce qui est admis.

3.1 Température thermodynamique

On considère le système (Σ) représenté sur la figure 38.1. Deux gaz (Σ_1) et (Σ_2) occupent deux compartiments de volume V_1 et V_2 , entourés par une enceinte adiabatique et indéformable. Les deux compartiments sont séparés par une cloison fixe diatherme (qui permet les transferts thermiques). Leurs températures thermodynamiques sont notées T_1 et T_2 et on fait l'hypothèse $T_1 > T_2$.

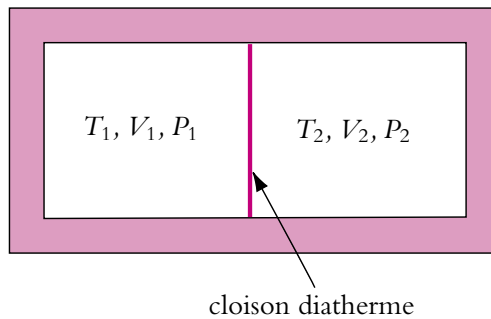


Figure 38.1 Mise à l'équilibre thermique.

On applique le premier principe à (Σ) entre deux états d'équilibre voisins. Sachant que la paroi est indéformable et adiabatique, (Σ) ne reçoit ni travail ni transfert thermique soit :

$$dU = \delta W + \delta Q = 0 \quad (38.2)$$

Spontanément, le gaz du compartiment de gauche se refroidit et celui du compartiment de droite se réchauffe jusqu'à l'égalité des températures. Le deuxième principe doit donc décrire l'évolution du système dans ce sens. Le deuxième principe doit aussi permettre de démontrer que le système Σ_1 perd de l'énergie. Entre deux états voisins, la variation élémentaire d'entropie de l'ensemble du système (Σ) sera :

$$dS = \delta S_{\text{créée}} + \delta S_{\text{éch}} = \delta S_{\text{créée}} > 0 \quad (38.3)$$

puisque l'évolution est adiabatique, $\delta S_{\text{éch}} = 0$ et puisqu'elle est aussi irréversible (transfert thermique spontané du corps chaud vers le corps froid), $\delta S_{\text{créée}} > 0$. On

utilise maintenant les propriétés d'extensivité de U et de S . Les expressions (38.2) et (38.3) donnent :

$$dS_1 + dS_2 > 0 \quad \text{et} \quad dU_1 + dU_2 = 0 \quad (38.4)$$

Lors des évolutions précédentes, les systèmes évoluaient à volume constante donc l'entropie et l'énergie interne ne dépendaient que de la température.

On définit la *température thermodynamique* comme le rapport de proportionnalité entre dU et dS lors des évolutions précédentes (c'est-à-dire à volume constant). Ainsi, on aura pour les sous-systèmes Σ_1 et Σ_2 :

$$\begin{aligned} dU_1 &= T_1 dS_1 \\ dU_2 &= T_2 dS_2 \end{aligned} \quad (38.5)$$

Maintenant qu'on possède tous les éléments, on peut démontrer que $dU_1 < 0$, ce qui traduit le fait que Σ_1 perd de l'énergie. L'inégalité entre les températures donne :

$$\frac{1}{T_1} < \frac{1}{T_2} \Rightarrow \frac{dS_1}{dU_1} < \frac{dS_2}{dU_2}$$

En utilisant la relation entre dU_1 et dU_2 (38.4), on obtient :

$$\frac{dS_1}{dU_1} < -\frac{dS_2}{dU_1} \Rightarrow \frac{dS_1 + dS_2}{dU_1} < 0$$

Or d'après (38.4), $dS_1 + dS_2 > 0$, donc :

$$dU_1 < 0 \quad \text{et} \quad dU_2 = -dU_1 > 0 \quad (38.6)$$

Le deuxième principe et la température thermodynamique ainsi définie montrent que l'échange d'énergie se fait spontanément du corps le plus chaud vers le corps le plus froid, ce qui est compatible avec l'expérience. Ceci valide l'énoncé du deuxième principe et la définition (38.5).

On a vu qu'à l'équilibre, S est maximale et que $dS = 0$. Avec les équations (38.5) et (38.6), on a :

$$dS = dS_1 + dS_2 = \frac{dU_1}{T_1} + \frac{dU_2}{T_2} = dU_1 \left(\frac{1}{T_1} - \frac{1}{T_2} \right) \quad (38.7)$$

Puisque $dS = 0$ à l'équilibre, on en déduit que $T_1 = T_2$.

En fait, la définition de la température thermodynamique est plus générale. L'évolution précédente a eu lieu à volume constant, ce qu'il faut préciser dans la définition. On définit donc la température thermodynamique par :

$$\frac{1}{T} = \left(\frac{\partial S}{\partial U} \right)_V \quad (38.8)$$

On admet que la température thermodynamique s'identifie à la température absolue.

3.2 Pression thermodynamique

Dans le paragraphe précédent, on a supposé la paroi fixe. Que se passe-t-il si l'on libère la paroi, une fois les températures équilibrées ? On sait d'expérience que la cloison va se déplacer jusqu'à ce que les pressions s'équilibrent de part et d'autre. Encore une fois, le deuxième principe doit permettre de décrire cette évolution.

De manière générale, pour un système d'énergie interne U , d'entropie S et de volume V évoluant entre deux états voisins, comme l'entropie S est une fonction de l'énergie interne U et du volume V , $S(U, V)$, on peut écrire :

$$dS = \left(\frac{\partial S}{\partial U} \right)_V dU + \left(\frac{\partial S}{\partial V} \right)_U dV \quad (38.9)$$

Avec la définition de la température, le premier terme de la somme est égal à $\frac{1}{T}$. On utilise une analyse dimensionnelle pour justifier la définition de la pression. Les termes d'une somme doivent être de même dimension, (ici le rapport d'une énergie sur une température exprimé en $\text{J} \cdot \text{K}^{-1}$), on en déduit que :

$$\left(\frac{\partial S}{\partial V} \right)_U \text{ s'exprime en } \frac{\text{J} \cdot \text{K}^{-1}}{\text{m}^3} \quad (38.10)$$

Or le produit PV est homogène à une énergie (Cf. travail des forces de pression) d'où :

$$\left(\frac{\partial S}{\partial V} \right)_U \text{ s'exprime en } \frac{\text{J} \cdot \text{K}^{-1}}{\text{m}^3} = \frac{\text{Pa} \cdot \text{m}^3 \cdot \text{K}^{-1}}{\text{m}^3} = \text{Pa} \cdot \text{K}^{-1} \quad (38.11)$$

Ainsi, la deuxième dérivée partielle est homogène au rapport d'une pression sur une température. Ceci permet de définir la pression thermodynamique par :

$$\frac{P}{T} = \left(\frac{\partial S}{\partial V} \right)_U \quad (38.12)$$

On admet que cette pression s'identifie aux deux précédentes : pression cinétique et pression de l'équation d'état du gaz parfait.

On revient alors à l'expérience de la figure 38.1 en supposant que les températures de Σ_1 et Σ_2 sont égales mais que les pressions P_1 et P_2 (initialement $P_1 > P_2$) sont différentes et qu'on lâche la paroi. Le raisonnement est similaire à celui effectué au paragraphe précédent. L'évolution est adiabatique et la paroi extérieure est indéformable (V est constant) donc entre deux états voisins :

$$dU = dU_1 + dU_2 = \delta W + \delta Q = 0 \quad \text{et} \quad dV = 0 = dV_1 + dV_2 \quad (38.13)$$

L'évolution est adiabatique irréversible donc entre deux états voisins :

$$dS = dS_1 + dS_2 > 0 \quad (38.14)$$

Avec les définitions de T et P , l'expression précédente peut s'écrire :

$$\frac{dU_1}{T_1} + \frac{P_1}{T_1}dV_1 + \frac{dU_2}{T_2} + \frac{P_2}{T_2}dV_2 > 0 \quad (38.15)$$

Si on note T la température commune, en regroupant les termes et en utilisant les différentes relations précédentes, on déduit :

$$\frac{dU_1 + dU_2}{T} + \frac{P_1 dV_1 + P_2 dV_2}{T} > 0 \Rightarrow \frac{P_1 - P_2}{T} dV_1 > 0 \quad (38.16)$$

Ainsi puisque par hypothèse, $P_1 > P_2$, l'inégalité impose que $dV_1 > 0$ donc le volume de la partie gauche augmente alors qu'en contrepartie le volume de droite diminue. **L'évolution décrite par le deuxième principe est donc conforme à l'expérience.**

Quel est l'état final ? L'entropie S est maximale à l'équilibre donc $dS = 0$. De l'expression précédente de dS on tire :

$$\frac{P_1 - P_2}{T} dV_1 = 0 \quad (38.17)$$

Puisque le volume V_1 n'est pas constant, à l'équilibre : $P_1 = P_2$.

3.3 Identité thermodynamique

On peut maintenant écrire la relation reliant dS et dU pour une évolution quelconque d'un système fermé :

$$dS = \frac{dU}{T} + \frac{P}{T}dV \quad \text{ou} \quad dU = TdS - PdV \quad (38.18)$$

L'expression (38.18) portant le nom *d'identité thermodynamique* est très importante et très utilisée. C'est la définition de la température T et de la pression P , elle est donc valable dès que les termes qui la composent sont définis. On peut établir une seconde identité thermodynamique faisant intervenir l'enthalpie H et la pression P , intéressante à utiliser en particulier dans le cas d'une transformation isobare.

$$\begin{aligned} dH = dU + PdV + VdP &\Leftrightarrow dH = TdS + VdP \\ &\Leftrightarrow dS = \frac{dH}{T} - \frac{V}{T}dP \end{aligned} \quad (38.19)$$

► **Remarque :** La connaissance de la fonction entropie $S(U, V)$ ou $S(H, P)$ d'un gaz permet, grâce à l'identité thermodynamique, de déterminer son équation d'état et l'expression de son énergie interne. $S(U, V)$ et $S(H, P)$ sont des **fonctions caractéristiques** du gaz.

3.4 Cas d'une transformation réversible

Dans le cas d'une transformation réversible pour laquelle les seules forces intervenant sont les forces de pression, le travail s'écrit $\delta W_{\text{rév}} = -PdV$. Alors, d'après le premier principe : $dU = -PdV + \delta Q_{\text{rév}}$. Après comparaison avec l'identité thermodynamique (38.18), on obtient :

$$\delta Q_{\text{rév}} = TdS \quad (38.20)$$

On a vu que, dans le cas général, on pouvait écrire : $dS = \delta S_{\text{créée}} + \delta S_{\text{éch}}$. Or ici la transformation est réversible donc $\delta S_{\text{créée}} = 0$. On en déduit, dans le cas d'une transformation réversible, une expression de dS qui, dans ce cas, est égal $\delta S_{\text{éch}}$ soit :

$$dS = \delta S_{\text{éch}} = \frac{\delta Q_{\text{rév}}}{T} \quad (38.21)$$

La relation (38.20) est donc, heureusement, compatible avec l'écriture du second principe (relation (38.21)).

On déduit de l'équation précédente (38.21) que si on trouve une transformation réversible allant du même état initial au même état final (S étant une fonction d'état, sa variation ne dépend que de l'état final et de l'état initial), on pourra calculer facilement la variation d'entropie par $dS = \frac{\delta Q_{\text{rév}}}{T}$. Cependant, si la transformation étudiée n'est pas réversible, on n'aura pas égalité entre $\delta S_{\text{éch}}$ et dS , car $\delta S_{\text{créée}} \neq 0$.

4. Exemples de calculs d'entropie

4.1 Entropie d'un gaz parfait

Le système considéré est un gaz parfait. Entre deux états d'équilibre voisins, l'expression de dU (relation (35.18) ou première loi de Joule) et l'identité thermodynamique (38.18) donnent :

$$dU = TdS - PdV = C_V dT = \frac{nR}{\gamma - 1} dT$$

soit :

$$dS = \frac{nR}{\gamma - 1} \frac{dT}{T} + \frac{P}{T} dV = \frac{nR}{\gamma - 1} \frac{dT}{T} + \frac{nR}{V} dV \quad (38.22)$$

ce qui donne par intégration :

$$S = \frac{nR}{\gamma - 1} \ln \left(\frac{T}{T_0} \right) + nR \ln \left(\frac{V}{V_0} \right) + S_0 \quad (38.23)$$

où S_0 est l'entropie du gaz pour la température T_0 et le volume V_0 . On peut regrouper les deux termes :

$$S = \frac{nR}{\gamma - 1} \ln \left(\frac{TV^{\gamma-1}}{T_0 V_0^{\gamma-1}} \right) + S_0 \quad (38.24)$$

► **Remarque** : La constante S_0 dépend du nombre de moles du système.

On peut aussi exprimer l'entropie en fonction des variables T et P . Dans ce cas, on part de l'identité thermodynamique sous la forme (relation 38.18) et de la seconde loi de Joule :

$$dH = TdS + VdP = C_P dT = \frac{nR\gamma}{\gamma - 1} dT$$

soit :

$$dS = \frac{nR\gamma}{\gamma - 1} \frac{dT}{T} - \frac{V}{T} dP = \frac{nR\gamma}{\gamma - 1} \frac{dT}{T} - \frac{nR}{P} dP \quad (38.25)$$

ce qui donne par intégration :

$$S = \frac{nR}{\gamma - 1} \ln \left(\frac{T}{T_0} \right) - nR \ln \left(\frac{P}{P_0} \right) + S_0 \quad (38.26)$$

où S_0 est l'entropie du gaz pour la température T_0 et la pression P_0 (c'est la même que précédemment). En regroupant les deux termes :

$$S = \frac{nR}{\gamma - 1} \ln \left(\frac{P^{1-\gamma} T^\gamma}{P_0^{1-\gamma} T_0^\gamma} \right) + S_0 \quad (38.27)$$

Si on utilise l'équation du gaz parfait, on peut exprimer S en fonction des variables P et V .

$$S = \frac{nR}{\gamma - 1} \ln \left(\frac{PV^\gamma}{P_0 V_0^\gamma} \right) + S_0 \quad (38.28)$$

Dans le cas d'une transformation adiabatique réversible, la variation d'entropie est nulle. Si on utilise les trois équations précédentes entre l'état initial (indice i) et l'état final (indice f), en calculant la différence $S_f - S_i = 0$, on retrouve les équations de Laplace (37.9).

4.2 Variation d'entropie d'un gaz parfait lors de quelques transformations simples

On peut, dans certains cas, calculer la variation d'entropie sans utiliser les formules générales (38.24), (38.28) et (38.27). Il suffit de repartir de l'identité thermodynamique et d'intégrer entre l'état initial (indice i) et l'état final (indice f).

a) Transformation isochore

Dans ce cas, $dV = 0$ soit :

$$dS = \frac{nR}{\gamma - 1} \frac{dT}{T}$$

et en intégrant :

$$\Delta S = \frac{nR}{\gamma - 1} \ln \left(\frac{T_f}{T_i} \right)$$

b) Transformation isotherme

Dans ce cas, $dT = 0$ et $dU = 0$ puisque le gaz est parfait.

$$dS = nR \frac{dV}{V}$$

et en intégrant :

$$\Delta S = nR \ln \left(\frac{V_f}{V_i} \right)$$

c) Transformation isobare

Dans ce cas, il vaut mieux utiliser l'expression en fonction de P et H (équation (38.19)), sachant que $dP = 0$:

$$dS = \frac{dH}{T} = \frac{nR\gamma}{\gamma - 1} \frac{dT}{T}$$

soit en intégrant :

$$\Delta S = \frac{nR\gamma}{\gamma - 1} \ln \left(\frac{T_f}{T_i} \right)$$

Ces expressions peuvent bien sûr être retrouvées à partir des expressions générales.

4.3 Refroidissement ou échauffement d'un corps

On étudie l'évolution d'un corps solide de capacité thermique C constante, initialement à la température T_1 et plongé dans une très grande quantité d'eau à température T_2 . La pression extérieure est constante et égale à la pression atmosphérique. La quantité d'eau étant très importante, on suppose que l'eau reste à température constante T_2 . Dans ce cas, la température finale du solide est celle de l'eau, soit T_2 .

On calcule la variation d'entropie du corps. Pour un corps solide, $dH = dU = CdT$. Dans ce cas, en appliquant l'identité thermodynamique avec $dV = 0$ (ou $dP = 0$), puisque le corps est solide, on peut écrire :

$$dS = \frac{dU}{T} = \frac{dH}{T} = C \frac{dT}{T} \quad (38.29)$$

En intégrant entre l'état initial et l'état final, on obtient, puisque la capacité thermique C est constante :

$$\Delta S = \int_{T_1}^{T_2} C \frac{dT}{T} = C \ln \left(\frac{T_2}{T_1} \right) \quad (38.30)$$

La variation d'entropie donnée par l'expression (38.30) peut être positive ou négative suivant les valeurs des températures. Cela n'est pas en contradiction avec le deuxième principe puisque l'évolution n'est pas adiabatique.



Il faut retenir l'expression de la variation d'entropie élémentaire (38.29) pour un corps solide qui est aussi valable pour un liquide incompressible.

4.4 Notion de thermostat ou source de chaleur

Dans l'exemple précédent, on a supposé que la température de l'eau ne variait pas. On dit que l'eau s'est comportée comme un *thermostat*. L'hypothèse est licite car la capacité thermique du corps solide a été considérée petite devant celle de l'eau. Par exemple, un baigneur nageant dans la mer ne va pas en modifier la température. Cette notion de thermostat est très importante en thermodynamique et il faut en donner une définition plus précise.

Un *thermostat* (ou source de chaleur) est un système de capacité thermique infinie, n'échangeant pas de travail mais uniquement du transfert thermique avec l'extérieur. Il est entièrement caractérisé par sa température.

Cette définition correspond à un thermostat idéal qui, grâce à sa capacité thermique infinie, reste à température constante. En effet, puisqu'il est caractérisé par sa température, son énergie interne ne dépend que de T , soit : $\Delta U = C\Delta T$. Or C étant infinie, on en déduit que, comme ΔU doit rester finie, ΔT est nulle.

Dans la pratique, deux types de thermostats peuvent être utilisés. Le plus simple est d'avoir à disposition un système de capacité thermique très grande devant le système étudié comme dans le paragraphe précédent. On peut citer en particulier les exemples suivants :

- pour les moteurs de voitures, on utilise l'atmosphère comme source froide (l'échange se fait au niveau du radiateur) ;
- pour les réfrigérateurs, on utilise l'atmosphère comme source chaude (échangeur derrière le réfrigérateur) ;
- pour les pompes à chaleur, on peut utiliser comme source froide, suivant le type :
 - l'atmosphère (pompe air-air),
 - une grande quantité d'eau comme un lac ou un fleuve (pompe air-eau) ;

- pour les centrales électriques (thermiques ou nucléaires), on utilise soit l'atmosphère, l'échange se faisant grâce à des immenses tours de refroidissement, soit un fleuve ;
- dans certaines expériences en chimie, on utilise l'eau du robinet avec une colonne de refroidissement.

On reviendra plus en détail sur la notion de source chaude et froide dans le chapitre sur les machines thermiques, mais ce qu'on peut constater dès maintenant, c'est que **toutes ces sources sont gratuites**.

Un autre type de thermostat très utilisé dans les laboratoires est le bain thermostaté. Il s'agit le plus généralement d'eau dont on mesure la température régulièrement et qu'on maintient à température constante grâce à une résistance chauffante.

Ainsi en pratique, on utilise en général comme thermostats des systèmes évoluant à pression ou volume constants qui sont bien caractérisés par leur température uniquement.

On applique les principes de la thermodynamique à un thermostat noté Σ . On note Q le transfert thermique reçu par Σ (le travail est nul par définition) entre l'état initial et l'état final. Si le système évolue à pression constante, le premier principe donne : $\Delta H = Q$ et si c'est à volume constant : $\Delta U = Q$. En fait, comme pour les solides ou les liquides incompressibles, on peut en confondre U et H et écrire :

$$\Delta U \simeq \Delta H = Q \quad (38.31)$$

Pour calculer leur variation d'entropie, il suffit d'appliquer l'une des expressions de l'identité thermodynamique (38.18) ou (38.19), dans lesquelles seul le terme en T n'est pas nul ($dP = 0$ ou $dV = 0$). Soit pour un thermostat :

$$dS = \frac{dU}{T} = \frac{dH}{T} \Rightarrow \Delta S = \int_{E_i}^{E_f} \frac{dH}{T} = \frac{1}{T} \int_{E_i}^{E_f} dH \Rightarrow \Delta S = \frac{\Delta H}{T} = \frac{Q}{T} \quad (38.32)$$

Attention, il existe des sources de chaleur imparfaites, pour lesquelles il est nécessaire de tenir compte de la valeur finie de la capacité thermique C . C'est le cas par exemple si, au lieu d'utiliser un lac, on utilise une tonne d'eau. Alors, on ne peut plus considérer que la température est constante et les calculs de ΔU , ΔH et ΔS se font comme pour un solide, soit pour une évolution de la température T_i à la température T_f :

$$\begin{aligned} \Delta U \simeq \Delta H &= Q = C (T_f - T_i) \\ \Delta S &= \int_{T_i}^{T_f} \frac{dT}{T} = C \ln \left(\frac{T_f}{T_i} \right) \end{aligned} \quad (38.33)$$

4.5 Analyse du système global thermostat - corps en contact

On examine à nouveau le cas du paragraphe 4.3 d'un corps (S) à la température T_1 en contact avec un thermostat (Σ) à la température T_2 . Le seul échange qui existe est le transfert thermique entre l'eau et le corps : l'ensemble est donc isolé de l'extérieur comme cela est schématisé sur la figure 38.2. On applique les deux principes aux différents systèmes.

$$\text{Corps (S)} \left| \begin{array}{l} 1^{\text{er}} \text{ principe : } \Delta H_S = C(T_2 - T_1) = Q \\ 2^{\text{e}} \text{ principe : } \Delta S_S = C \ln \left(\frac{T_2}{T_1} \right) \end{array} \right. \quad (38.34)$$



Le transfert thermique est orienté vers (S) et non vers (Σ), d'où le signe – dans l'écriture des deux principes :

$$\text{Thermostat } (\Sigma) \left| \begin{array}{l} 1^{\text{er}} \text{ principe : } \Delta H_{\Sigma} = -Q \\ 2^{\text{e}} \text{ principe : } \Delta S_{\Sigma} = \frac{\Delta H_{\Sigma}}{T_2} = \frac{-Q}{T_2} \end{array} \right. \quad (38.35)$$

$$\text{Ensemble } (\Sigma, S) \left| \begin{array}{l} 1^{\text{er}} \text{ principe : } \Delta H_T = \Delta H_{\Sigma} + \Delta H_S = 0 \\ 2^{\text{e}} \text{ principe : } \Delta S_T = \Delta S_{\Sigma} + \Delta S_S = \frac{-Q}{T_2} + C \ln \left(\frac{T_2}{T_1} \right) \end{array} \right. \quad (38.36)$$

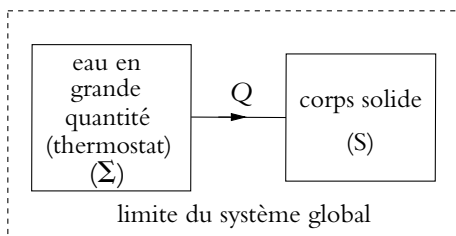


Figure 38.2 Différents systèmes.

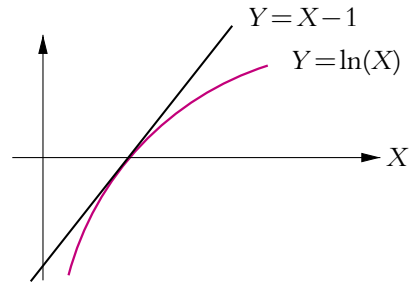


Figure 38.3 La courbe $Y = \ln(X)$ est au dessous de la courbe $Y = X - 1$.

Des équations précédentes, on tire l'expression de la variation de l'entropie totale en fonction des températures :

$$\Delta S_T = C \left(\ln \left(\frac{T_2}{T_1} \right) + \frac{T_1}{T_2} - 1 \right) = C \left(X - 1 - \ln(X) \right) \quad \text{avec } X = \frac{T_1}{T_2} \quad (38.37)$$

Sur la figure 38.3 sont tracées les courbes $X - 1$ et $\ln(X)$, et on s'aperçoit que, sauf pour $X = 1$ ($T_2 = T_1$), la quantité $X - 1 - \ln(X)$ est toujours strictement positive.

Ainsi, pour l'ensemble du système, la variation d'entropie est strictement positive ($\Delta S_T > 0$), **ce qui est en accord avec le deuxième principe puisque l'évolution de l'ensemble (S, Σ) est adiabatique irréversible.**

Par abus de langage, on appelle souvent l'entropie de l'ensemble du système évoluant adiabatiquement : *variation d'entropie de l'Univers*.

On peut calculer pour le corps solide, l'entropie échangée et l'entropie créée. Par définition, l'entropie échangée l'est avec le thermostat, c'est d'ailleurs l'opposé de la variation d'entropie du thermostat : en effet, d'après la relation (38.1),

$$S_{\text{éch}} = \frac{Q}{T_2}$$

La variation d'entropie du thermostat est :

$$\Delta S_{\Sigma} = \frac{-Q}{T_2}$$

d'après ce qui précède, d'où :

$$S_{\text{éch}} = -\Delta S_{\Sigma}$$

On en déduit l'entropie créée :

$$S_{\text{créée}} = \Delta S_S - S_{\text{éch}} = C \left(\ln \left(\frac{T_2}{T_1} \right) + \frac{T_1}{T_2} - 1 \right)$$

et on trouve que **l'entropie créée est égale à la variation d'entropie totale et est bien strictement positive.**



Trois remarques importantes :

- La seule manière de calculer l'entropie créée est de calculer la variation d'entropie du corps puis de retrancher l'entropie échangée.
- L'entropie échangée est $\int \frac{\delta Q}{T_{\Sigma}} = \frac{\int \delta Q}{T_2} = \frac{Q}{T_2}$ et c'est la température du thermostat qui intervient.
- L'entropie créée l'est par le corps. Un thermostat parfait évolue de manière réversible.

On s'intéresse maintenant au cas où les températures T_2 et T_1 sont proches. On note ΔT la différence $T_2 - T_1$. On calcule $S_{\text{créée}}$ en effectuant un développement limité sachant que $|\Delta T| \ll T_1$ et $|\Delta T| \ll T_2$.

$$\begin{aligned} S_{\text{créée}} &= C \left(-\ln \left(\frac{T_2 - \Delta T}{T_2} \right) + \frac{T_2 - \Delta T}{T_2} - 1 \right) \\ &\simeq C \left(\left(\frac{\Delta T}{T_2} + \frac{1}{2} \left(\frac{\Delta T}{T_2} \right)^2 \right) - \frac{\Delta T}{T_2} \right) \end{aligned}$$

Après simplification, l'équation précédente donne :

$$S_{\text{créée}} = \frac{C}{2} \left(\frac{\Delta T}{T_2} \right)^2 \quad (38.38)$$

$$\text{D'autre part, } Q = C(T_2 - T_1) = C\Delta T \quad (38.39)$$

D'après les expressions (38.38) et (38.39), on remarque que le transfert thermique est d'ordre 1 en $\frac{\Delta T}{T_2}$, alors que la création d'entropie est d'ordre 2. Donc lorsque les températures sont proches, $S_{\text{créée}}$ tend plus vite vers 0 que Q et la transformation est quasi réversible, car $S_{\text{créée}} \ll \frac{Q}{T_2}$. Ainsi, pour rendre la transformation réversible, il faut mettre le corps en contact successivement avec une multitude de thermostats dont les températures sont réparties entre T_1 et T_2 .

On verra d'ailleurs que le moyen d'effectuer des échanges réversibles avec des corps à températures constantes utilise des changements d'état par exemple liquide-vapeur.

En conclusion, on peut dire que la cause de l'irréversibilité de la transformation précédente est la différence de température entre le corps et le thermostat.

4.6 Compression isotherme d'un gaz parfait

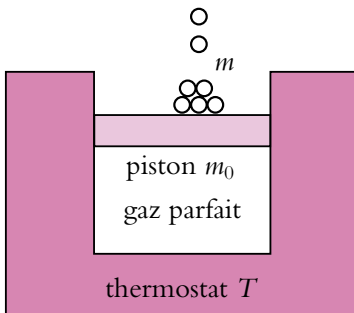


Figure 38.4 Cas réversible.

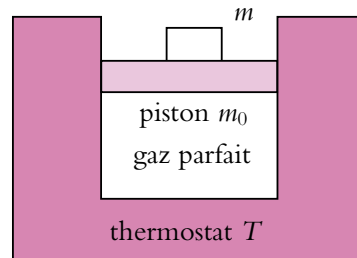


Figure 38.5 Cas irréversible.

On va étudier successivement la compression isotherme d'un gaz parfait effectuée de manière réversible puis irréversible. On suppose que le gaz est dans un cylindre fermé par un piston adiabatique de section s et que les parois diathermes du cylindre sont en contact avec un thermostat (Σ) à la température T . Le piston a une masse m_0 et est surmonté de vide (s'il y avait une atmosphère au-dessus du piston, il faudrait ajouter la pression de l'atmosphère à la pression extérieure). On ajoute une masse m sur le piston soit progressivement (cas réversible de la figure 38.4), soit brusquement (cas irréversible de la figure 38.5). Dans les deux cas, on veut calculer l'entropie créée $S_{\text{créée}}$. Le système (S) est constitué du gaz, du piston et de l'enceinte. En supposant que les capacités thermiques de l'enceinte et du piston sont très petites, on pourra négliger leur variation d'énergie interne devant celle du gaz.

a) Cas réversible

Pour le système (S) :

$$\text{État initial (1)} \left| \begin{array}{l} V_1 \\ P_1 = \frac{m_0 g}{s} \\ T \end{array} \right| \quad \left| \right| \quad \text{État final (2)} \left| \begin{array}{l} V_2 \\ P_2 = \frac{(m_0 + m)g}{s} \\ T \end{array} \right|$$

Le piston part du repos et est de nouveau au repos dans l'état final donc sa variation d'énergie cinétique est nulle : $\Delta E_c = 0$. On applique le premier principe à (S), en négligeant ΔU_{piston} et $\Delta U_{\text{enceinte}}$ devant ΔU_{gaz} , ce qui donne :

$$\Delta U_S = W + Q \Leftrightarrow \Delta U_{\text{gaz}} + \Delta U_{\text{piston}} + \Delta U_{\text{enceinte}} + \Delta E_c = W + Q$$

soit :

$$\Delta U_{\text{gaz}} = W + Q$$

Puisque le gaz parfait subit une évolution isotherme $\Delta U_{\text{gaz}} = C_V \Delta T = 0$ et puisqu'elle est quasi-statique, le travail et le transfert thermique sont donnés par l'expression (37.3) soit :

$$Q = -W = -nRT \ln \left(\frac{P_2}{P_1} \right) = -nRT \ln \left(1 + \frac{m}{m_0} \right) \quad (38.40)$$

Pour calculer la variation d'entropie du gaz, on utilise les résultats du paragraphe 4.1 ou on repart de l'identité thermodynamique ou encore puisque la transformation est réversible, on applique $dS_{\text{gaz}} = \frac{\delta Q_{\text{rév}}}{T}$. On choisit cette dernière méthode. Dans ce cas, sachant que la température est constante, on peut écrire :

$$\Delta S_{\text{gaz}} = \int dS_{\text{gaz}} = \frac{1}{T} \int \delta Q = \frac{Q}{T} \Rightarrow \Delta S_{\text{gaz}} = -nR \ln \left(1 + \frac{m}{m_0} \right) \quad (38.41)$$

Comme pour l'énergie interne, on néglige les variations d'entropie de l'enceinte et du piston devant celle du gaz.

L'entropie échangée est égale à :

$$S_{\text{éch}} = \frac{Q}{T} = nR \ln \left(1 + \frac{m}{m_0} \right) \quad (38.42)$$

On remarque que c'est l'opposé de la variation d'entropie du thermostat (attention au signe car Q est orienté vers (S)) :

$$\Delta S_{\Sigma} = -\frac{Q}{T} = -nR \ln \left(1 + \frac{m}{m_0} \right) \quad (38.43)$$

Des expressions précédentes, on tire l'entropie créée :

$$S_{\text{créée}} = \Delta S_{\text{gaz}} - S_{\text{éch}} = \frac{Q}{T} - \frac{Q}{T} = 0 \quad (38.44)$$

L'entropie créée est nulle, ce qui est en accord avec le deuxième principe, puisque le système (S) subit une transformation réversible.

b) Cas irréversible

Puisque les pressions initiale et finale et la température sont les mêmes que dans le cas réversible, les volumes sont aussi les mêmes : le gaz est donc dans le même état initial et dans le même état final qu'au cas précédent. Puisque S et U sont des fonctions d'état, elles ne dépendent que des états extrêmes, ΔU_{gaz} et ΔS_{gaz} sont identiques à ceux calculés précédemment.

On va calculer l'entropie créée par le système global et donc le transfert thermique Q . Comme précédemment $W = -Q$ puisque $\Delta U_{\text{gaz}} = 0$. Cependant, la transformation n'étant pas quasi-statique, on ne peut appliquer l'expression (37.3). Dès que la masse m est posée, on peut dire que la force extérieure s'exerçant sur (S) et qui travaille est constante et égale à $(m + m_0)g$. Donc la pression extérieure exercée sur (S) est égale à la pression finale P_2 . On peut alors calculer le travail reçu par (S) soit :

$$W = - \int P_{\text{ext}} dV = -P_{\text{ext}} \int_{V_1}^{V_2} dV = -P_{\text{ext}}(V_2 - V_1)$$

Ce qui peut s'écrire :

$$W = \frac{(m + m_0)g}{s} V_1 \left(1 - \frac{V_2}{V_1}\right) \Leftrightarrow W = \frac{(m + m_0)g}{s} \frac{nRT}{P_1} \left(1 - \frac{P_1}{P_2}\right)$$

Il suffit alors de remplacer P_1 et P_2 par leurs expressions en fonction des masses et on obtient après simplification :

$$W = nRT \frac{m}{m_0} - Q$$

On obtient donc pour l'entropie échangée :

$$S_{\text{éch}} = \frac{Q}{T_{\Sigma}} = \frac{Q}{T} = -nR \frac{m}{m_0} \quad (38.45)$$

On peut alors calculer l'entropie créée :

$$S_{\text{créée}} = \Delta S_{\text{gaz}} - S_{\text{éch}} = nR \left(-\ln \left(1 + \frac{m}{m_0}\right) + \frac{m}{m_0} \right) \quad (38.46)$$

On trouve dans ce cas que l'entropie créée est strictement positive, ce qui confirme que la transformation est irréversible.

Ainsi, bien que le gaz se trouve dans les mêmes états dans les deux cas, les transferts thermiques et les travaux sont différents et l'entropie créée est aussi différente car l'état final du thermostat n'est pas le même. La deuxième transformation est irréversible car, en raison du déplacement rapide du piston, il apparaît des inhomogénéités de densité importantes.

Elle n'est pas quasi-statique, en effet la pression du gaz n'est pas égale à la pression extérieure, elle ne peut donc pas être réversible.

4.7 Détentes de gaz parfait

Dans le chapitre sur les applications du premier principe on a étudié deux types de détente : la détente de Joule - Gay Lussac et celle de Joule - Thomson, toutes deux adiabatiques. D'après leurs descriptions, on a conclu qu'elles étaient irréversibles ; on va le confirmer en effectuant un bilan entropique dans les deux cas pour un gaz parfait. Ces deux détente étant adiabatiques, l'entropie échangée est nulle donc $\Delta S_{\text{gaz}} = S_{\text{créée}}$, il suffit donc de calculer la variation d'entropie du gaz.

Il suffit d'appliquer les expressions de l'entropie d'un gaz parfait (38.23) pour la détente de Joule - Gay Lussac et (38.26) pour la détente de Joule - Thomson. On a démontré que dans les deux cas la température reste constante et on la note T .

a) Détente de Joule - Gay Lussac :

Soit V_1 le volume initial et V_2 le volume final. Alors l'expression (38.23) donne :

$$\Delta S_{\text{gaz}} = \frac{nR}{\gamma - 1} \ln \left(\frac{T}{T} \right) + nR \ln \left(\frac{V_2}{V_1} \right) = nR \ln \left(\frac{V_2}{V_1} \right) \quad (38.47)$$

or $V_2 > V_1$ donc $\Delta S_{\text{gaz}} > 0$.

La détente de Joule - Gay Lussac, adiabatique, est bien irréversible. Lors de l'ouverture du robinet, de fortes inhomogénéités de densité apparaissent au sein du gaz. Ce sont elles qui sont à l'origine de l'irréversibilité de la transformation.

► **Remarque :** Pour n'importe quel gaz, si on connaît l'expression $S(U, V)$, on peut calculer la variation d'entropie lors d'une détente de Joule - Gay Lussac.

b) Détente de Joule - Thomson :

Soit P_1 la pression initiale et P_2 la pression finale. Alors l'expression (38.26) donne :

$$\Delta S_{\text{gaz}} = \frac{nR\gamma}{\gamma - 1} \ln \left(\frac{T}{T} \right) - nR \ln \left(\frac{P_2}{P_1} \right) = nR \ln \left(\frac{P_1}{P_2} \right) \quad (38.48)$$

or $P_2 < P_1$ donc $\Delta S_{\text{gaz}} > 0$. La détente de Joule - Thomson, adiabatique, est bien irréversible. L'origine de l'irréversibilité de la transformation est la diffusion de particules entre les deux zones de concentrations différentes.

► **Remarque :** Pour n'importe quel gaz, si on connaît l'expression $S(H, P)$, on peut calculer la variation d'entropie lors d'une détente de Joule - Thomson.

5. Relation de Clausius

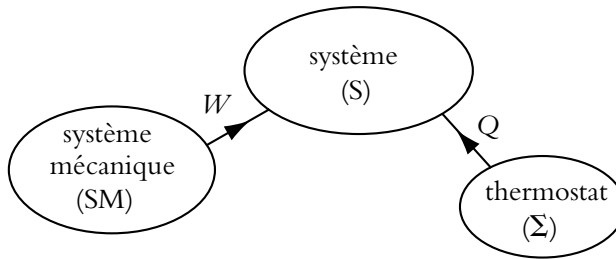


Figure 38.6 Convention d'orientation du travail et du transfert thermique.

On utilise les propriétés de l'entropie, pour établir une inégalité entre le transfert thermique reçu par un système et sa variation d'entropie. On étudie un système (S) en contact avec un thermostat (Σ) à température T_Σ et un système mécanique (SM) n'échangeant que du travail (Cf. figure 38.6). (S) ne reçoit du transfert thermique que de (Σ) et du travail que de (SM). On peut écrire :

$$S_{\text{créée}} \geq 0 \quad (38.49)$$

Comme le transfert thermique ne s'effectue qu'entre (S) et (Σ), on a d'après les orientations :

$$S_{\text{éch}} = \frac{Q}{T_\Sigma} = \frac{Q}{T}$$

et puisque :

$$\Delta S_S = S_{\text{éch}} + S_{\text{créée}}$$

on en déduit :

$$\Delta S_S \geq \frac{Q}{T_\Sigma} \quad (38.50)$$

Cette inégalité qui va être généralisée à plusieurs thermostats porte le nom d'*inégalité de Clausius*. Il y a **égalité dans le cas d'une évolution réversible**. Une évolution où le système est en contact avec un seul thermostat comme ce qui vient d'être décrit, est appelée évolution *monotherme*. On parle d'évolution *ditherme* dans le cas de deux thermostats. On se place dans le cas de N thermostats (Σ_i), de température T_i . Avec les mêmes conventions d'orientation que dans le cas monotherme, (S) reçoit un transfert thermique Q_i de chacun des thermostats. On peut de nouveau écrire

$$S_{\text{créée}} \geq 0 \quad (38.51)$$

Or ici le transfert thermique se fait avec les N thermostats, donc l'entropie échangée s'écrit :

$$S_{\text{éch}} = \sum_{i=1}^N \frac{Q_i}{T_i}$$

d'où :

$$\Delta S_S \geq \sum_{i=1}^N \frac{Q_i}{T_i} \quad (38.52)$$

Cette équation constitue l'inégalité de Clausius.

On aurait pu d'ailleurs la choisir comme énoncé du deuxième principe, c'est ce qui a été fait dans le passé.

Dans le cas où les sources de chaleur ne sont pas à température constante, il faut tenir compte de l'évolution de leur température. Les expressions précédentes deviennent alors :

$$S_{\text{éch}} = \sum_{i=1}^N \int \frac{\delta Q_i}{T_i}$$

et

$$\Delta S_S \geq \sum_{i=1}^N \int \frac{\delta Q_i}{T_i}$$

A. Applications directes du cours

1. Bilan entropique

Un morceau de fer de 2 kg, chauffé à blanc (à la température de 880 K) est jeté dans un lac à 5° C. Quelle est l'entropie créée ? On donne la capacité calorifique massique du fer : $c_{\text{fer}} = 4600 \text{ kJ.kg}^{-1}.\text{K}^{-1}$. Quelle est la cause de cette création d'entropie ?

2. Mise à l'équilibre

Les 2 compartiments d'un cylindre isolé de volume V contiennent initialement un gaz parfait monoatomique à la même température $T = 300 \text{ K}$ sous les pressions $P_1 = 1 \text{ bar}$ et $P_2 = 2 \text{ bar}$; les volumes initiaux sont égaux à $\frac{V}{2} = 1 \text{ L}$. La paroi diatherme qui sépare les deux compartiments est rendue libre.

1. Déterminer l'état final.
2. Calculer l'entropie créée.

3. Transformations polytropiques d'un gaz parfait diatomique

Soit un gaz parfait diatomique tel que le produit PV^k soit constant. On note γ le rapport des capacités calorifiques à pression et à volume constants, n le nombre de moles et R la constante des gaz parfaits.

1. Établir l'expression différentielle de l'entropie et la mettre sous la forme $f(\gamma, n, R, k) \frac{dT}{T}$.
2. En déduire la variation d'entropie au cours de la transformation entre deux états A et B .
3. En supposant la transformation quasistatique, calculer le travail W et le transfert thermique Q .
4. A quelle situation correspondent les cas :
 1. $k = 1$?
 2. $k = 0$?
 3. k infini ?
 4. $k = \gamma$?

4. Évolution d'une enceinte, d'après ENAC 2004

Un récipient parfaitement isolé thermiquement est composé de deux compartiments séparés par une paroi isolante thermiquement. Initialement chaque compartiment contient un gaz parfait dont les capacités thermiques molaires à volume ou à pression constantes sont notées c_v ou c_p . On note n_i , P_i et T_i le nombre de moles, la pression et la température du compartiment i . On enlève la paroi séparant les deux compartiments.

1. Calculer la température finale T_f du mélange qu'on considérera comme un gaz parfait.
1. $T_f = \frac{T_1 + T_2}{n_1 + n_2}$,

$$2. T_f = \frac{n_1 T_1 + n_2 T_2}{2},$$

$$3. T_f = \frac{n_1 T_1 + n_2 T_2}{n_1 + n_2},$$

$$4. T_f = \frac{T_1 + T_2}{2}.$$

2. Exprimer la pression finale P_f du mélange.

$$1. P_f = P_1 P_2 \frac{n_1 T_1 + n_2 T_2}{n_1 T_1 P_1 + n_2 T_2 P_2},$$

$$2. P_f = P_1 P_2 \frac{n_1 T_1 + n_2 T_2}{T_1 P_2 + T_2 P_1},$$

$$3. P_f = P_1 P_2 \frac{n_1 T_1 + n_2 T_2}{T_1 P_1 + T_2 P_2},$$

$$4. P_f = P_1 P_2 \frac{n_1 T_1 + n_2 T_2}{n_1 T_1 P_2 + n_2 T_2 P_1}.$$

3. Calculer la pression finale $P_{f,i}$ de chaque gaz i en le considérant comme seul (on parle alors de pression partielle).

$$1. P_{f,1} = \frac{n_2}{n_1 + n_2} P_f \text{ et } P_{f,2} = \frac{n_1}{n_1 + n_2} P_f,$$

$$2. P_{f,1} = \frac{n_1}{n_1 + n_2} P_f \text{ et } P_{f,2} = \frac{n_2}{n_1 + n_2} P_f,$$

$$3. P_{f,1} = \frac{n_1 n_2}{n_1 + n_2} P_f = P_{f,2},$$

$$4. P_{f,1} = P_{f,2} = P_f.$$

4. Calculer la variation d'entropie du système constitué par les deux gaz pour $n_1 = n_2 = 1$ mole. On remarque qu'il faut utiliser les pressions partielles.

$$1. \Delta S = c_p \ln \frac{T_f^2}{T_1 T_2} + R \ln \frac{P_1 P_2}{P_f^2} + 2R \ln 2,$$

$$2. \Delta S = c_p \ln \frac{T_1 T_2}{T_f^2} + R \ln \frac{P_f^2}{P_1 P_2} + 2R \ln 2,$$

$$3. \Delta S = c_p \ln \frac{T_1 T_2}{T_f^2} + R \ln \frac{P_1 P_2}{P_f^2} - 2R \ln 2,$$

$$4. \Delta S = c_p \ln \frac{T_f^2}{T_1 T_2} + R \ln \frac{P_1 P_2}{P_f^2}.$$

5. Que devient ce résultat si $P_1 = P_2 = P_0$ et $T_1 = T_2 = T_0$?

$$1. \Delta S = 0,$$

$$2. \Delta S = -2R \ln 2,$$

$$3. \Delta S = 2R \ln 4,$$

$$4. \Delta S = 2R \ln 2.$$

5. Bilan entropique du chauffage d'un morceau de cuivre

- On chauffe à l'aide d'un thermostat à la température $T_S = 600 \text{ K}$ une mole de cuivre solide pour faire passer sa température de 293 K à 320 K . On néglige la variation du volume et on admet que la capacité thermique est $3 R$. Effectuer un bilan entropique de ce chauffage. Conclure.
- On fait passer un courant électrique de $5,0 \text{ A}$ dans un conducteur ohmique de 44Ω pendant 1 h . La puissance fournie par effet Joule sert à chauffer une enceinte à la température de 293 K et fait passer la température de la résistance chauffante de 293 K à 313 K . On admet que la capacité thermique de la résistance chauffante est de $5,0 \text{ J.K}^{-1}$. Effectuer le bilan entropique de l'opération. Conclure.

6. Fonction caractéristique

On trouve expérimentalement l'expression de l'entropie d'une mole de dioxyde de carbone : $S(U, V) = S_0 + c_V \ln \frac{U + \frac{a}{V}}{U_0 + \frac{a}{V_0}} + R \ln \frac{V - b}{V_0 - b}$ avec $c_V = 28,5 \text{ J.mol}^{-1}.\text{K}^{-1}$, $a = 0,37 \text{ J.m}^3.\text{mol}^{-2}$ et $b = 4,30.10^{-5} \text{ J.mol}^{-1}.\text{m}^3$.

- Exprimer la différentielle de l'entropie.
- En déduire l'équation d'état du gaz ainsi que l'expression de son énergie interne.
- Que deviennent ces expressions pour n moles de gaz ?
- On fait subir à $5,0 \text{ L}$ de ce gaz une détente de Joule-Gay Lussac pour atteindre un volume final de 10 L . Déterminer la variation de température au cours de cette transformation.
- Calculer la variation d'entropie associée.
- Étudier en la justifiant la dépendance de ces variations avec le nombre de moles.
- Que peut-on conclure de l'ensemble de cette étude ?

7. Évolution isotherme et évolution adiabatique

On considère une masse de 100 g d'eau dans laquelle plonge un conducteur de résistance $R = 20 \Omega$. Cette dernière est parcourue par un courant de 10 A pendant 1 s . On note (S) le système formé de l'eau et de la résistance. On donne :

- masse du conducteur : $m_c = 19 \text{ g}$;
- capacité thermique massique du conducteur : $c_c = 0,42 \text{ J.g}^{-1}.\text{K}^{-1}$;
- capacité thermique massique de l'eau : $c_e = 4,18 \text{ J.g}^{-1}.\text{K}^{-1}$.

- La température de l'ensemble est maintenue constante et égale à 20°C . Quelle est la variation d'entropie de (S) ? Quelle est l'entropie créée? Quelle est la cause de la création d'entropie?
- Le même courant passe dans le conducteur pendant la même durée mais maintenant (S) est isolé thermiquement. Calculer la variation d'entropie de (S) et l'entropie créée. Quelle est la cause de la création d'entropie?

8. Cycle en diagramme entropique

1. On considère un diagramme, appelé *diagramme entropique*, dans lequel on trace la température en ordonnées et l'entropie en abscisse. Dans ce diagramme, on étudie une transformation cyclique réversible. Montrer que le transfert thermique reçu par le système au cours du cycle est représenté par l'aire du cycle. Dans quel sens doit être parcouru ce cycle pour qu'il soit moteur ?
2. On considère dans le plan (P, V) les isothermes $T_0, 2T_0, 3T_0, \dots, nT_0$, tracées pour une mole du gaz parfait ainsi que les isentropiques $S_0, 2S_0, 3S_0, \dots, nS_0$. On obtient ainsi un maillage dans ce plan (P, V) . Montrer que les aires de tous les "carreaux" de ce maillage sont égales.

B. Exercices et problèmes

1. Mise en contact avec des thermostats

1. Une masse $m = 1$ kg d'eau liquide à $T_0 = 273$ K est mise en contact avec un thermostat à $T_f = 300$ K. La capacité thermique massique de l'eau est $c = 4,18 \cdot 10^3$ J.kg⁻¹.K⁻¹. Calculer l'entropie créée. Quelle est la cause de la création d'entropie ?
2. On suppose maintenant que l'eau est d'abord mise en contact avec un premier thermostat à la température $T_1 = 285$ K jusqu'à ce qu'elle atteigne cette température, puis elle est mise en contact avec le thermostat à T_f . Calculer de nouveau l'entropie créée. Pourquoi trouve-t-on une valeur inférieure à celle de la question précédente ?
3. On opère maintenant en N étapes : l'eau, initialement à $T_0 = 273$ K, est mise en contact avec N thermostat de température T_i vérifiant $\frac{T_i}{T_{i-1}} = a = \text{cste}$ et $T_N = T_f = 300$ K. Calculer l'entropie créée. Que devient-elle quand N tend vers l'infini ? Pourquoi ?

2. Sens d'un cycle monotherme

Une mole de gaz parfait ($\gamma = 1,4$) subit la succession de transformations suivante :

- détente isotherme de $P_A = 2$ bar et $T_A = 300$ K jusqu'à $P_B = 1$ bar, en restant en contact avec un thermostat à $T_T = 300$ K ;
 - évolution isobare jusqu'à $V_C = 20,5$ L toujours en restant en contact avec le thermostat à T_T ;
 - compression adiabatique réversible jusqu'à l'état A .
1. Représenter ce cycle en diagramme (P, V) . S'agit-il d'un cycle moteur ou récepteur ?
 2. Déterminer l'entropie créée entre A et B .
 3. Calculer la température en C , le travail W_{BC} et le transfert thermique Q_{BC} reçus par le gaz au cours de la transformation BC . En déduire l'entropie échangée avec le thermostat ainsi que l'entropie créée.
 4. Calculer la valeur numérique de l'entropie créée au cours d'un cycle. Le cycle proposé est-il réalisable ? Le cycle inverse l'est-il ?

3. Entrée de matière dans un récipient

On considère un récipient vide cylindrique de volume V_1 , constitué de parois adiabatiques. On perce un trou de manière à ce que l'air ambiant (P_0 , T_0) y pénètre de façon adiabatique (transformation rapide). On appelle V_0 le volume initialement occupé par l'air qui entre dans le récipient.

1. Calculer la température finale de l'air du récipient.
2. Déterminer l'entropie créée. Quelle est la cause de la création d'entropie ?

4. Détente irréversible d'un gaz parfait

Soit le dispositif de la figure ci-contre. Les parois et le piston sont adiabatiques. La paroi interne est fixe et diatherme (elle permet les échanges thermiques). Elle est percée d'un trou fermé par une fenêtre amovible.

La pression extérieure est $P_0 = 1$ bar. Initialement le volume A est rempli d'un gaz parfait ($P_0 = 1$ bar, $T_0 = 300$ K, $n = 1$ mol) et le volume B est vide.

Le rapport des capacités thermiques du gaz γ vaut 1, 4.

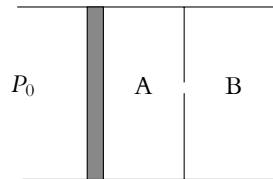


Figure 38.7

1. On ouvre la fenêtre. Décrire qualitativement ce qui se passe suivant la taille de l'enceinte B . En déduire l'existence d'un volume critique de B : V_C que l'on ne demande pas de calculer.
2. On suppose $V_B < V_C$.
 - a) On appelle V_1 le volume final occupé par le gaz. Déterminer le travail reçu par le gaz.
 - b) Déterminer l'état final du gaz (P_1 , V_1 , T_1) en fonction de P_0 , V_A et V_B .
 - c) Calculer l'entropie créée. Conclure. Quelle est la cause de la création d'entropie ?
 - d) Déterminer V_C . Effectuer l'application numérique.
3. Reprendre les questions 2 dans le cas $V_B > V_C$
4. On suppose maintenant que seul le piston est adiabatique, et que le dispositif est maintenu à T_0 par un thermostat. On appelle V'_C le nouveau volume critique.
 - a) Déterminez le nouvel état final pour $V_B < V'_C$.
 - b) Déterminez le nouveau volume critique V'_C .
 - c) Calculer l'entropie créée quand $V_B < V'_C$.

5. Étude de quelques transformations du gaz parfait, d'après CCP TSI 2006

On étudie différentes transformations de n moles d'un gaz parfait. On note P sa pression, V son volume, T sa température, c_V sa capacité calorifique molaire à volume constant, c_P sa capacité calorifique molaire à pression constante et $\gamma = \frac{c_P}{c_V}$.

1. Donner l'équation des gaz parfaits en précisant les unités des variables et les valeurs des constantes.

2. Établir la relation de Mayer entre les capacités calorifiques molaires à volume et à pression constants.
3. En déduire les expressions de c_p et c_V en fonction de γ et R la constante des gaz parfaits.
4. Rappeler la différentielle de l'énergie interne pour les gaz parfaits. En déduire celle de l'entropie en utilisant les variables T et V puis les variables T et P .
5. On enferme le gaz dans une enceinte diathermane (permettant les transferts thermiques) dont une paroi de masse négligeable est mobile verticalement sans frottement. Le milieu extérieur se comporte comme un thermostat de température T_1 . Initialement le gaz est caractérisé par une pression P_1 , un volume V_1 et une température T_1 . On suppose que la paroi est dans un premier temps bloquée. On la débloque et on la déplace de manière quasistatique jusqu'à ce que le volume V'_1 offert au gaz soit égal à $2V_1$ et on la rebloque.
Déterminer la pression P'_1 du gaz à l'état final.
6. Déterminer le travail W_1 mis en jeu au cours de cette transformation en fonction de n , R et T_1 .
7. Calculer la variation d'énergie interne au cours de la transformation et en déduire le transfert thermique reçu par le gaz.
8. Donner l'expression de l'entropie échangée avec le milieu extérieur.
9. Établir la variation d'entropie au cours de la transformation.
10. En déduire l'expression de l'entropie créée. Que peut-on en conclure ?
11. On suppose maintenant que les parois sont athermanes (ne permettant pas les échanges thermiques). Le récipient est divisé en deux compartiments de volume identique, le gaz se trouvant dans l'un d'eux et le deuxième étant vide. On enlève la séparation sans fournir de travail.
Établir l'expression de la variation d'énergie interne au cours de la transformation.
12. Déterminer les valeurs de la température T'_1 et de la pression P'_1 à l'équilibre.
13. Si on n'avait pas considéré un gaz parfait, quelle différence aurait-on observée ?
14. Déterminer la variation d'entropie au cours de l'évolution.
15. Quelle est l'entropie échangée ?
16. En déduire l'entropie créée. Que peut-on en conclure ?
17. Comparer ces deux transformations.

6. Thermodynamique d'un baromètre, d'après ENSTIM 2006

1. On assimile un baromètre à un corps solide de capacité calorifique C initialement à la température T_1 . On le plonge dans un lac à la température T_0 constante et on attend l'équilibre thermique. Exprimer la variation d'entropie du baromètre ΔS_B en fonction de C et $x = \frac{T_1}{T_0}$.
2. Exprimer la variation d'entropie du lac ΔS_L en fonction des mêmes variables.
3. En déduire la variation d'entropie ΔS de l'ensemble baromètre et lac.

4. Montrer par une étude graphique que ΔS est toujours positive.
5. On reprend le baromètre à sa température initiale T_1 et le lac à la température T_0 constante inférieure à T_1 . Ces deux éléments sont assimilables à deux sources de chaleur qu'on utilise pour fabriquer un moteur effectuant des cycles réversibles. On note δQ_B le transfert thermique entre la machine et le baromètre, δQ_L le transfert thermique entre la machine et le lac et δW le travail fourni par la machine au cours d'un cycle. Au cours d'un cycle, la température du baromètre passe de la valeur T comprise entre T_1 et T_0 à la température $T + dT$.
Quelle égalité peut-on écrire entre δW , δQ_B et δQ_L ?
6. Quelle inégalité peut-on écrire entre δQ_B , δQ_L , T et T_0 ?
7. Que deviennent ces relations dans le cas réversible ?
8. Le moteur s'arrête de fonctionner quand la température du baromètre atteint la valeur T_0 . Exprimer les transferts thermiques Q_B et Q_L durant la totalité du fonctionnement du moteur en fonction de T_1 et T_0 .
9. Définir le rendement de ce moteur en fonction de W et Q_B .
10. Donner son expression en fonction de C , T_1 et T_0 .

Interprétation statistique de l'entropie (PCSI)

39

1. Microétat et macroétat

Soit le jeu suivant : on possède trois boules semblables et trois récipients de tailles différentes. Le jeu consiste à lancer les boules dans les récipients sachant qu'une boule dans le plus petit (noté P) rapporte 3 points, une boule dans le moyen (noté M) rapporte 2 points et une boule dans le grand (noté G) rapporte 1 point. Toute boule à l'extérieur ne rapporte aucun point. Quelles possibilités y a-t-il pour que le joueur gagne 3 points ? Les solutions sont représentées sur la figure 39.1.

- Solution (1) : 2 boules à l'extérieur et 1 boule dans P ;
- Solution (2) : 3 boules dans G ;
- Solution (3) : 1 boule à l'extérieur, une boule dans G et une boule dans M.

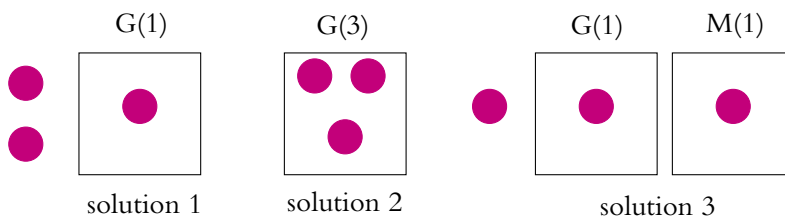


Figure 39.1 Les différentes solutions de gagner trois points.

On voit donc qu'à un seul total de points (ici 3 points) correspondent 3 possibilités de répartition des boules. Maintenant si on remplace les boules par des particules et les boîtes par des niveaux d'énergie (0, E , $2E$, $3E$) d'un système, on voit que pour mesurer une énergie égale à $3E$, il y aura aussi trois répartitions possibles à l'état microscopique :

- Solution (1) : 2 particules à 0 et 1 particule à $3E$;

- Solution (2) : 3 particules à E ;
- Solution (3) : 1 particule à $2E$, une particule à 0 et une particule à E .

Ce que mesure l'observateur au niveau macroscopique, c'est l'état macroscopique (appelé *macroétat*) qui est donc composé de plusieurs états microscopiques appelés *microétats*. On donne une définition plus précise de ces deux grandeurs.

Un macroétat est un état du système décrit par un ensemble de variables macroscopiques, ces variables étant en petit nombre.

Un microétat correspond à la répartition des particules à l'échelle microscopique (position, vitesse, énergie...).

Un macroétat peut correspondre à plusieurs microétats. On note Ω le nombre de microétats possibles pour un macroétat.

Hypothèse : **Tous les états microscopiques possibles, correspondant à la même énergie macroscopique (totale), sont équiprobables.** Cette hypothèse est dite *hypothèse microcanonique*.

2. État macroscopique le plus probable

On s'intéresse au système représenté sur la figure 39.2, constitué de N particules, confinées dans une enceinte indéformable, de volume V , isolée de l'extérieur par une paroi adiabatique. On a disposé au milieu de l'enceinte une cloison ouverte au centre, qui délimite deux volumes égaux. L'observateur s'intéresse au nombre de particules par unité de volume en un point de l'enceinte.

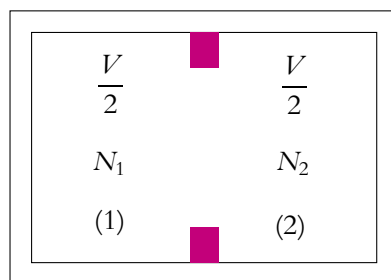


Figure 39.2 Répartition du gaz entre les deux compartiments.

On note N_1 le nombre de particules dans le volume de gauche et $N_2 = N - N_1$ celui dans le volume de droite. On cherche à déterminer N_1 et N_2 . Combien y a-t-il de possibilités de choisir N_1 particules parmi N ? On est ramené à un problème d'analyse combinatoire, dont le résultat est $C_N^{N_1}$. Le nombre de possibilités n_{N_1} d'avoir

N_1 particules dans le volume gauche est donc :

$$n_{N_1} = C_N^{N_1} = \frac{N!}{N_1!(N - N_1)!} \quad (39.1)$$

Il s'agit maintenant de déterminer la probabilité d'avoir N_1 particules à gauche. Pour cela il faut déterminer le nombre total de possibilités d'avoir des particules à gauche, c'est-à-dire N_1 variant de 0 à N . Le nombre total n_t de possibilités est donc la somme de tous les n_{N_1} . On retrouve un résultat mathématique connu, le binôme de Newton, soit :

$$n_t = \sum_{N_1=0}^N n_{N_1} = \sum_{N_1=0}^N C_N^{N_1} = 2^N \quad (39.2)$$

La probabilité $p(N_1)$ de trouver N_1 particules dans le compartiment gauche est donc :

$$p(N_1) = \frac{C_N^{N_1}}{2^N} \quad (39.3)$$

Les figures 39.3 et 39.4 représentent l'allure des probabilités en fonction de N_1 pour $N = 10$ et pour N très grand (de l'ordre du nombre d'Avogadro par exemple). Dans le cas où N est très grand, les points du tracé étant très rapprochés, la courbe semble continue. On remarque que dans les deux cas **la probabilité est la plus forte pour** $N_1 = N_2 = \frac{N}{2}$, résultat attendu, on note $p_{\max}$ la probabilité maximale. D'autre part, si N est très grand, la courbe apparaît plus étroite et plus piquée que pour $N = 10$.

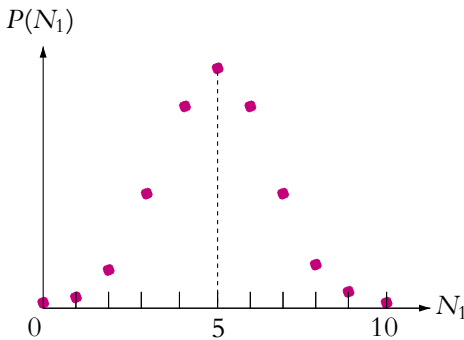


Figure 39.3 Cas $N = 10$.

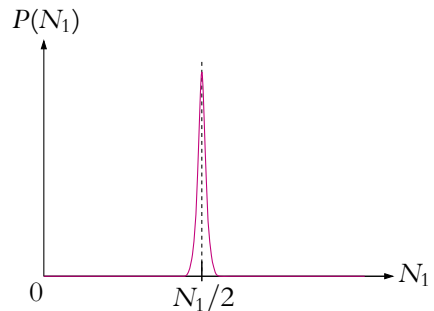


Figure 39.4 Cas N très grand.

On quantifie ceci en évaluant la largeur de la courbe à mi-hauteur $\Delta_{1/2} = n_2 - n_1$, où n_1 et n_2 sont les valeurs de N_1 telles que $p(N_1) = \frac{p_{\max}}{2}$. On ne développera pas ici le calcul sachant seulement qu'il faut effectuer un développement de l'expression (39.3) au voisinage de $\frac{N}{2}$, en utilisant la formule d'approximation de Stirling valable pour

N grand :

$$N! = \sqrt{2\pi N} N^N \exp(-N) \quad (39.4)$$

On montre alors qu'au voisinage de $\frac{N}{2}$ la probabilité peut s'écrire :

$$p(N_1) = p_{\max} \exp \left(-2 \frac{\left(N_1 - \frac{N}{2}\right)^2}{N} \right) \quad (39.5)$$

On détermine ainsi les valeurs de n_1 et de n_2 :

$$\exp \left(-2 \frac{\left(N_1 - \frac{N}{2}\right)^2}{N} \right) = \frac{1}{2} \Leftrightarrow \left(N_1 - \frac{N}{2}\right)^2 = \frac{N}{2} \ln(2) \quad (39.6)$$

On en déduit les valeurs de n_1 , n_2 et $\Delta_{1/2}$:

$$n_{1/2} = \frac{N}{2} \pm \frac{1}{\sqrt{2}} \sqrt{N \ln(2)} \quad \text{et} \quad \Delta_{1/2} = \sqrt{2 \ln 2 N} \quad (39.7)$$

La largeur représente donc le double de la variation de N_1 lorsque la probabilité a été divisée par deux. Il est plus intéressant de se rendre compte si ce nombre est important par rapport au nombre total de particules N . C'est donc la variation relative de N_1 , $\Delta_r = \frac{\Delta_{1/2}}{N}$ qui est la grandeur intéressante :

$$\Delta_r = \frac{\Delta_{1/2}}{N} = \frac{\sqrt{2 \ln(2)}}{\sqrt{N}} \quad (39.8)$$

On s'aperçoit donc que la variation relative de N_1 est inversement proportionnelle à $\sqrt{N}$, ce qui confirme que plus le nombre N de particules est important, plus la courbe est étroite. Par exemple pour $N = 10^{24}$, Δ_r est de l'ordre de 10^{-12} , soit $10^{-10} \%$, ce qui est très faible. Cela signifie que la probabilité décroît très rapidement quand N_1 s'éloigne de la valeur $\frac{N}{2}$ et que l'on peut considérer que seules les possibilités de répartition des particules avec N_1 voisin de $\frac{N}{2}$ ont une probabilité non nulle.

En quelque sorte, si $N = 10$, il n'est pas impossible d'observer 2 particules à gauche et 8 = $N - 2$ à droite, mais si $N = 10^{24}$, on ne pourra pas observer 2 particules à gauche et $10^{24} - 2$ particules à droite.

Les résultats précédents peuvent s'interpréter aussi en concluant que les fluctuations autour de la valeur moyenne $\frac{N}{2}$ sont d'autant plus faibles que la valeur de N est grande. On s'apercevra très peu d'une variation de 100 000 particules dans un compartiment si N vaut 10^{24} , par contre cela serait très visible si N valait 200 000.

En conclusion de ce paragraphe, on peut dire qu'on observera macroscopiquement $N_1 = N_2 = \frac{N}{2}$, car cela correspond au seul état n'ayant pas une probabilité d'apparition négligeable et que les variations de N_1 et N_2 sont inobservables macroscopiquement.

3. Entropie statistique

3.1 Définition et propriétés

On s'intéresse à l'expérience de la détente de Joule - Gay Lussac dans le cas particulier où le volume initialement vide est égal à la moitié du volume total. On note N le nombre de particules et les indices sont les mêmes que pour le paragraphe précédent. Dans l'état initial, on a donc $N_1 = N$ et $N_2 = 0$ et dans l'état final, d'après les conclusions du paragraphe précédent, on observe l'état $N_1 = N_2 = \frac{N}{2}$ avec quelques fluctuations (Cf. figure 39.5).

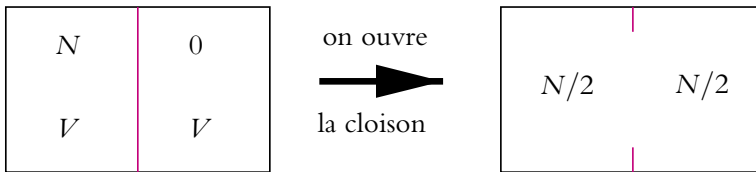


Figure 39.5 Détente de Joule - Gay Lussac.

Ce type de résultat a amené le physicien allemand Ludwig Boltzmann (1844-1906) à énoncer le principe suivant :

Pour un système isolé, l'état macroscopique observé est l'état macroscopique le plus probable, c'est-à-dire celui dont le nombre de microétats accessibles est maximal. On définit l'entropie statistique correspondant à un macroétat par :

$$S = k_B \ln \Omega \quad (39.9)$$

où Ω est le nombre de microétats correspondant au macroétat considéré.

La constante k_B porte le nom de *constante de Boltzmann*. Elle vaut $k_B = 1,38 \cdot 10^{-23} \text{ J.K}^{-1}$. Cet énoncé est effectivement en accord avec les résultats du paragraphe précédent sachant que dans l'hypothèse microcanonique tous les microétats sont équiprobables, le macroétat état dont la probabilité est la plus grande est bien celui formé du plus grand nombre de microétats.

On retrouve certains aspects du deuxième principe.

- Pour un système isolé, l'entropie est maximale à l'équilibre. En effet, d'après l'énoncé de Boltzmann, Ω est maximal donc l'entropie S aussi.
- L'entropie est extensive. En effet si le système (1) correspond à Ω_1 microétats et le système (2) à Ω_2 , lorsqu'on réunit les deux systèmes, le nombre total de microétats sera le produit $\Omega_1\Omega_2$, soit $S = S_1 + S_2$.

3.2 L'entropie statistique est-elle celle du deuxième principe ?

On a défini une fonction S appelée entropie grâce à l'énoncé de Boltzmann, mais est-ce que c'est la même fonction que celle définie dans l'énoncé du deuxième principe ? On va seulement vérifier la cohérence des deux définitions sur l'expérience de Joule - Gay Lussac décrite précédemment. Pour cela on calcule la variation d'entropie définie par l'expression (39.9) lors de la détente de la figure (39.5).

Dans l'état initial, $N_1 = N$ et $N_2 = 0$, soit $\Omega_i = C_N^0 = 1$ microétat et donc l'entropie initiale vaut $S_i = 0$.

Dans l'état final, d'après l'énoncé de Boltzmann, $N_1 = N_2 = \frac{N}{2}$, soit $\Omega_f = \frac{N!}{\left(\frac{N}{2}\right)!}$ et $S_f = k_B \ln \Omega_f$.

Donc :

$$S_f = k_B \ln N! - 2k_B \ln \left(\left(\frac{N}{2} \right)! \right)$$

On utilise la formule de Stirling (39.4) pour calculer S_f dans le cas d'un système thermodynamique pour lequel N est grand. Ainsi :

$$S_f = k_B \left(N \ln N - N + \frac{1}{2} \ln(2\pi N) - 2 \frac{N}{2} \ln \frac{N}{2} + 2 \frac{N}{2} - \ln(\pi N) \right)$$

On obtient après simplification :

$$S_f = k_B \left(N \ln 2 + \frac{1}{2} \ln 2 - \frac{1}{2} \ln(\pi N) \right) \quad (39.10)$$

On peut comparer l'ordre de grandeur des différents termes de l'équation précédente. Pour $N = 10^{24}$, le premier terme est de l'ordre de N , le second vaut environ 0,3, et le troisième terme vaut environ 30. On peut donc négliger les deux derniers termes et écrire $S_f = Nk_B \ln 2$. La variation d'entropie statistique lors de la détente vaut donc :

$$\Delta S = S_f - S_i = Nk_B \ln 2 \quad (39.11)$$

Dans le chapitre sur le deuxième principe, on a calculé la variation d'entropie pour cette détente (38.24) :

$$\Delta S = nR \ln \left(\frac{V_f}{V_i} \right) = Nk_B \ln \left(\frac{V_f}{V_i} \right)$$

Or ici le volume final est le double du volume initial donc $V_f = 2V_i$ soit $\Delta S = Nk_B \ln 2$. On trouve le même résultat avec les deux définitions de l'entropie. On n'a pas prouvé que les deux fonctions étaient semblables (ce qui est admis), mais on a bien vérifié la cohérence des deux définitions sur cet exemple.

Le deuxième principe permet d'affirmer que la détente de Joule - Gay Lussac est une transformation irréversible. Est-ce que l'énoncé de Boltzmann le permet aussi ? Pour cela, il faut calculer les probabilités correspondant à l'état initial et à l'état final. Pour l'état initial, puisque $\Omega_i = 1$, la probabilité $p_i = 2^{(-N)}$. Ainsi, pour $N = 10^{23}$ particules :

$$p_i = \frac{1}{2^{10^{23}}}$$

Le système aura passé statistiquement $\exp(-7.10^{22})$ seconde dans l'état initial. **Autant dire que l'on ne verra le système que dans l'état final le plus probable. La transformation ne pourra revenir vers l'état initial, c'est pour cela qu'elle est irréversible.**

4. Troisième principe de la thermodynamique

Dans le paragraphe précédent, on a vu que la variation d'entropie est la même pour les deux définitions dans le cas d'une détente de Joule - Gay Lussac. De plus, le deuxième principe ne fait pas allusion non plus à l'entropie mais à la variation d'entropie, et définit donc l'entropie à une constante près (dans une différence, les constantes disparaissent). Pour être sûr que les deux entropies se correspondent et pas seulement les variations d'entropie, il faut fixer la constante ; c'est ce que fait le troisième principe.

Troisième principe :

L'entropie d'un corps pur cristallisé parfait tend vers 0 J.K^{-1} lorsque la température tend vers 0 K .

D'après la définition de Boltzmann, on en déduit qu'à une température nulle, il n'y a qu'un seul microétat. En effet, l'énergie d'agitation thermique $k_B T$ est nulle, et la matière est complètement figée. Le troisième principe est indispensable pour établir les tables servant en thermochimie. Ces tables donnent les capacités thermiques des corps ainsi que les enthalpies et les entropies en fonction de l'état du corps à une

température donnée. On utilise des expressions qui seront vues dans le cours de chimie afin d'établir ces tables en prenant l'hypothèse d'une entropie nulle à 0 K.

5. Complément : entropie et information

À l'échelle macroscopique, l'observateur ne peut distinguer les différentes molécules dans l'exemple traité tout au long du chapitre. Il n'a accès qu'au nombre total de molécules présentes dans un compartiment : N_1 et N_2 .

- Si $N_1 = N$ et $N_2 = 0$, l'entropie S vaut 0 J.K^{-1} et il n'y a qu'un seul microétat. Toutes les particules sont dans le compartiment (1). Le système est parfaitement connu, **l'information est maximale**.
- Si $N_2 \neq 0$, on ignore dans quel compartiment se trouve une particule donnée puisqu'on ne connaît que le nombre de particules par compartiment. Il manque de l'information alors que S est supérieure à 0.

Ainsi S mesure l'information manquante d'un système, et plus S est grande moins on dispose d'information sur un système. Ainsi pour un même corps, si on appelle respectivement S_S , S_L et S_G l'entropie du corps sous forme solide, liquide et gazeuse, on aura l'inégalité suivante :

$$S_S < S_L < S_G \quad (39.12)$$

Le solide est en effet plus ordonné que le liquide, lui-même plus ordonné que le gaz : on connaît mieux la position des particules dans le solide que dans le gaz.

La notion d'entropie n'est pas utilisée qu'en thermodynamique. On la retrouve aussi en traitement du signal, en traitement d'image où une méthode appelée maximum d'entropie permet de reconstituer des images prises par satellites. Dans ce cas, tous les états n'ont pas forcément la même probabilité et on doit en tenir compte dans la définition de l'entropie. Dans le cas d'un signal, c'est la probabilité d'apparition d'un symbole ou d'un mot qui interviendra alors que dans une image, il s'agira de la probabilité d'apparition d'une couleur par exemple sur un pixel (pixel : plus petit détail d'une image).

Équilibre d'un corps pur sous plusieurs phases

40

1. Généralités

1.1 Solide, liquide, vapeur

Une *phase* est un système **homogène** ou un sous-système d'un système homogène. Dans une phase, **les paramètres intensifs sont uniformes**. Les corps purs peuvent exister sous trois phases : solide, liquide, vapeur. On utilise parfois le mot *état* à la place de phase.

En fait, il existe d'autres états qui seront présentés au paragraphe 1.3.

a) Phase vapeur ou gazeuse

- Un gaz occupe tout le volume qui lui est offert. Les interactions entre particules sont nulles (dans le cas d'un gaz parfait) ou faibles (force de Van der Waals).
- L'agitation thermique est importante.
- C'est un état très désordonné.



Une fumée n'est pas une vapeur, c'est une suspension de particules solides.

b) Phase liquide

- Les forces d'interaction sont plus grandes que pour les gaz, d'où un volume plus limité.
- La densité est plus grande que celle des gaz.
- Le volume varie peu avec la pression.
- Le liquide épouse la forme du récipient dans lequel il se trouve.

c) Phase solide

- C'est l'état le plus ordonné des trois.

- Les déplacements des particules sont limités autour de leur position d'équilibre.
- Il y a une forte cohésion.

Lorsque le corps évolue d'une phase à l'autre, on dit qu'il y a *changement de phase* ou *transition de phase*. Les différents changements sont définis sur la figure 40.1.

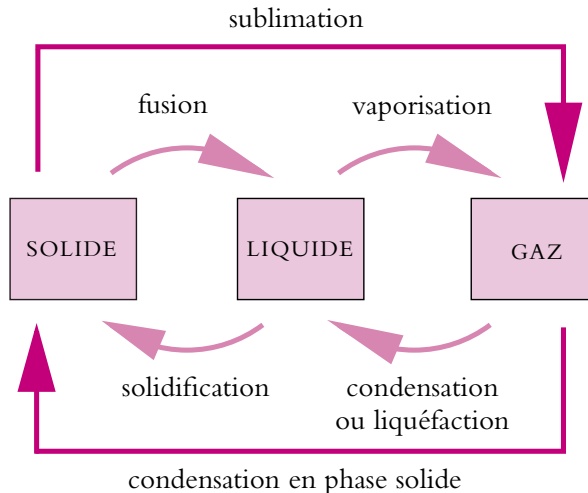


Figure 40.1 Le vocabulaire des changements de phase.

1.2 Utilisation des changements d'états

Il n'est pas question d'établir ici une liste exhaustive des phénomènes où l'on rencontre les changements d'états. On peut citer parmi les phénomènes naturels le cycle de l'eau qui fait alterner précipitation et évaporation. Dans la vie courante, certains appareils utilisent ces transformations tels les réfrigérateurs, congélateurs, climatiseurs, pompes à chaleur (voir le chapitre suivant sur les machines thermiques). Enfin on peut citer une application importante du point de vue industrielle : le « caloduc ». Ce système permet de transmettre une variation de température très rapidement d'un point à un autre, plus rapidement que n'importe quel métal. C'est un tuyau rempli d'un mélange liquide – vapeur dans les conditions critiques (Cf. paragraphe 2.2). Dès que la température augmente d'un côté du tuyau, une partie du liquide se vaporise et très rapidement l'effet se propage dans le tuyau assurant l'homogénéité de la température.

1.3 Variétés allotropiques

Pour certains corps, il existe plusieurs phases solides ou plusieurs phases liquides. Cela signifie que la structure microscopique n'est pas la même et que les propriétés physiques peuvent être différentes.

a) Exemples de corps ayant plusieurs phases solides

- L'eau : il existe 7 types de glace. La plus courante est la glace I. Lorsque la pression augmente (dans la banquise par exemple), la glace est comprimée et sa structure change. La température est aussi un facteur influant. On peut se convaincre d'ailleurs qu'il existe différentes glaces en observant un glaçon qui sort du congélateur ; l'extérieur est transparent et le centre blanc.
- Le soufre : S_α et S_β .
- Le fer : Fe_α , Fe_β , Fe_γ .

b) Exemple de corps ayant plusieurs phases liquides : l'hélium

En plus de l'hélium liquide ordinaire, il existe l'hélium II à des températures inférieures à 2 K et sous faible pression. L'hélium II est un superfluide (pas de frottement).

Il existe d'autres changements d'états, par exemple, pour les corps présentant des propriétés magnétiques comme le passage de l'état ferromagnétique (matière aimantée) à l'état paramagnétique (matière non aimantée). Ainsi, un aimant qu'on chauffe perd ses propriétés d'aimantation.

1.4 Retard au changement d'état et état métastable

Il arrive qu'un corps se trouve dans un état (une phase) même si la pression et la température correspondent théoriquement à une phase différente de celle dans laquelle il se trouve. Autrement dit le corps n'est pas dans l'état dans lequel il doit être théoriquement : on dit qu'il se trouve dans un *état métastable*. Cela signifie que la moindre perturbation va le faire passer à l'état stable thermodynamiquement.

Qui n'a pas vu des gouttelettes se former derrière les ailes d'un avion ou derrière une voiture quand l'air est très humide ? On dit dans ce cas que l'air est *sursaturé d'eau*, ce qui signifie qu'une grande (trop grande) quantité d'eau est présente à l'état vapeur et que le simple fait qu'un objet se déplace fait passer le surplus d'eau de la phase vapeur à la phase liquide. On appelle ce phénomène *un retard à la condensation*. Il était aussi utilisé pour visualiser des déplacements de particules (chambres dites de Wilson).

On prend maintenant l'exemple du phosphore : sous une pression de 1 bar, sa température de fusion est 44 °C. Si dans un tube à essai, on refroidit du phosphore liquide et très pur, doucement, on peut descendre jusqu'à une température de 35 °C sans observer de solidification. Dès qu'on frappe le tube, même légèrement, ou si on introduit une impureté, le phosphore se solidifie très rapidement et complètement. C'est ce qu'on appelle la *surfusion* du phosphore ou *retard à la solidification* observée avec d'autres corps tels le salol. Ce phénomène peut s'observer aussi avec une bouteille d'eau liquide retirée du congélateur qui, dans certains cas, se solidifie dès qu'elle est remuée. Cela signifie que la température de l'eau était inférieure à 0 °C même si elle était encore liquide.

On observe aussi des retards à l'ébullition : ne vous est-il jamais arrivé de retirer un bol d'eau du four à micro-onde et de voir, parce que vous l'avez bougé, des bulles apparaître et l'eau se mettre à bouillir ?

2. Diagramme d'équilibre

Les points représentatifs de l'état d'un corps décrit par les variables (P, V, T) définissent une surface appelée *surface des états*. Dans la pratique, il est cependant plus facile d'utiliser des diagrammes à deux dimensions qui sont les projections dans un plan de cette surface. On a l'habitude d'utiliser deux diagrammes : diagramme pression, température (P, T) et diagramme de Clapeyron (P, ν) où ν est le volume massique.

2.1 Diagramme pression, température : (P, T)

Un corps pur sous une phase peut être décrit par deux variables intensives, la pression et la température : on dit que le système est *divariant*. De façon empirique, on constate que, lorsque deux phases coexistent à l'équilibre, la pression et la température ne sont plus indépendantes. Dans ce cas, la pression est une fonction de la température, $P = f(T)$ et il suffit de connaître l'une des variables intensives pour décrire l'état du corps : on dit que le système est *monovariant*. Cette relation entre pression et température, qui sera justifiée en deuxième année, se traduit par une courbe d'équilibre (Cf. figure 40.2) dans le diagramme (P, T) séparant deux zones notées (1) et (2).

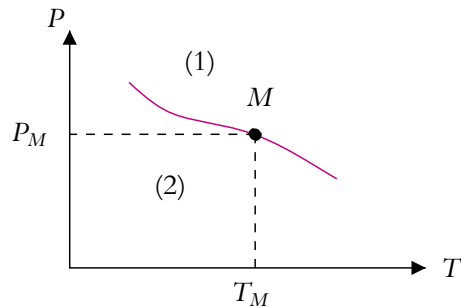


Figure 40.2 Diagramme d'équilibre entre deux phases quelconques.

- Dans les zones (1) et (2), seule l'une des phases est à l'équilibre et n'importe quel point représentatif est possible en choisissant la pression P et la température T tandis que sur la courbe, les deux phases coexistent et, si P est fixée, T l'est aussi (point M).
- Si trois phases coexistent, la pression P est reliée à la température T par deux fonctions (une par phase). Ainsi $P = f(T)$ et $P = g(T)$ et il n'existe qu'un seul couple (P, T) possible. On appelle *point triple* le point représentatif de ce couple dans le diagramme (P, T) . Ainsi, pour un corps décrit uniquement par la pression et la température, il ne peut y avoir coexistence de plus de trois phases.

Les figures 40.3 et 40.4 donnent l'allure des diagrammes solide, liquide, vapeur de la plupart des corps et le cas particulier de l'eau et du bismuth pour lesquels la courbe solide - liquide a une pente négative contrairement aux autres corps.

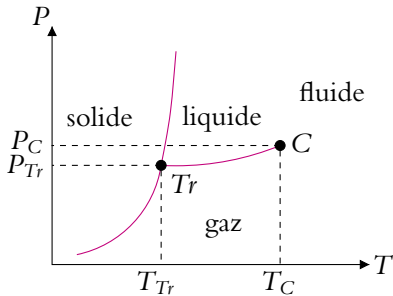


Figure 40.3 Diagramme (P, T) courant.

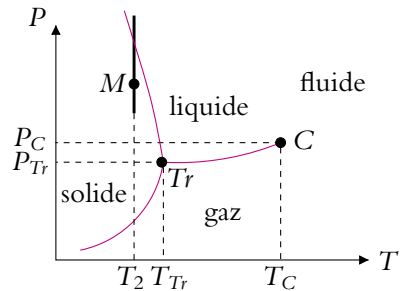


Figure 40.4 Diagramme (P, T) de l'eau et du bismuth.

Comment déterminer les zones de stabilité des trois phases ? Il suffit de se rappeler qu'à faible température et forte pression, on rencontre les corps plutôt sous forme solide et qu'au contraire à faible pression et forte température, le corps est gazeux.

Sur les diagrammes apparaissent deux points notés C et Tr . Le point Tr a déjà été évoqué, il s'agit du point triple pour lequel il y a coexistence de trois phases. Dans le cas de l'eau, la pression et la température du point triple sont : $P_{Tr} = 6,2 \cdot 10^{-3}$ bar et $T_{Tr} = 273,16$ K. En raison de la valeur de la pression, on ne rencontre pas couramment l'eau sous ses trois phases simultanément. Le point C quant à lui est appelé *point critique*. Il limite la courbe de vaporisation. Au-delà de ce point existe la zone de stabilité de l'état fluide (ou gaz hypercritique) pour laquelle le corps n'est ni liquide, ni gazeux mais dans un état intermédiaire appelé *état fluide*. Pour l'eau, la pression et la température du point critique sont : $P_C = 218$ bar et $T_C = 647$ °C. Lorsque la température notée T_1 est inférieure à T_C et que le gaz subit une compression isotherme, le point représentatif M de la transformation se déplace sur une droite verticale (trait continu sur la figure 40.5). Lorsque M atteint la courbe de condensation, des gouttelettes apparaissent dans le gaz qui commence à se liquéfier. Par contre, si on contourne le point critique entre les mêmes états que précédemment, par la courbe en pointillés de la figure 40.5, on n'assiste pas à l'apparition de gouttelettes mais à un passage continu du gaz au liquide.

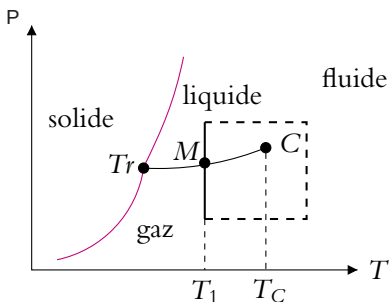


Figure 40.5 Liquéfaction directe et contournement du point critique.



Figure 40.6 Le fil de fer traverse la glace.

En raison de la pente négative de la courbe de solidification de l'eau, il est possible de réaliser l'expérience représentée sur la figure 40.6. On constate que le fil de fer s'enfonce dans la glace qui se referme après son passage. On peut interpréter l'expérience grâce au diagramme de l'eau. L'évolution isotherme à la température T_2 est représentée sur la figure 40.4 en trait vertical gras. Le point représentatif M de l'état initial est dans la zone solide. En raison des masses suspendues au fil, la pression sur la glace est très forte localement et le point représentatif passe dans la zone liquide, la glace fond et le fil s'enfonce. Ensuite la pression rediminue et la glace se reforme. Quand le fil a traversé la glace, le morceau de glace est intact, en un seul morceau. Cela ne peut pas se passer avec les corps dont les diagrammes sont ceux de la figure 40.3, car, dans ce cas, le point représentatif reste dans la zone solide.

2.2 Diagramme de Clapeyron pour l'équilibre liquide-vapeur

On se limite au cas de l'équilibre liquide-vapeur. En physique, on a l'habitude d'utiliser le diagramme de Clapeyron (Cf. figure 40.7) avec le volume massique ν (volume d'un kilogramme) en abscisse mais attention en chimie, c'est plutôt le volume molaire qui est utilisé. Le diagramme (P, T) est représenté en correspondance (Cf. figure 40.8).

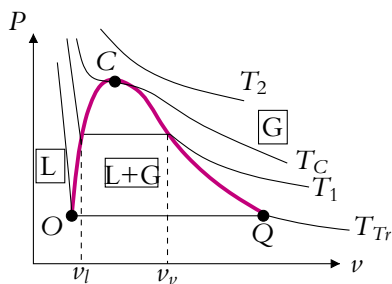


Figure 40.7 Diagramme de Clapeyron.

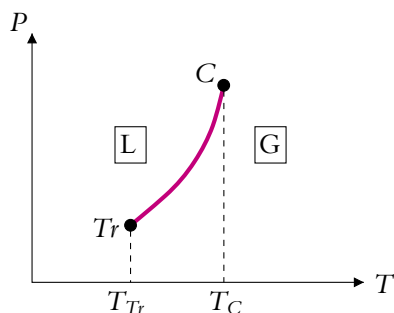


Figure 40.8 Diagramme (P, T) .

Il apparaît trois zones de stabilité notées L pour le liquide, G pour le gaz (ou vapeur) et L + G pour la zone d'équilibre entre la vapeur et le liquide. Cette partie correspond à une courbe projetée dans les coordonnées (P, T) . Cela signifie que **la pression et la température sont fixées mais que la composition du mélange peut varier**. Au-dessus du point critique se trouve la zone de l'état fluide. On a représenté quatre isothermes T_{Tr} , T_1 , T_C , et T_2 . L'isotherme T_{Tr} correspond à la température du point triple et est la limite du diagramme liquide - vapeur. L'isotherme T_C passe par le point critique et possède un point d'inflexion à tangente horizontale en C . On remarque la forme particulière des isothermes passant dans la zone L + G :

- une première partie dans la zone du liquide très pentue signifiant que le liquide étant peu compressible, la pression peut augmenter fortement tandis que le volume varie peu ;

- une partie horizontale dans la zone d'équilibre liquide - vapeur correspondant à un palier de changement d'état ;
- une partie dans la zone du liquide qui ressemble aux isothermes du gaz parfait.

Ces isothermes portent le nom d'*isothermes d'Andrews*. On a représenté la limite des paliers de changement d'état par une courbe en trait gras nommée *courbe de saturation*.

Entre le point O et le point C , la courbe de saturation est très pentue et porte le nom particulier de *courbe d'ébullition* et entre C et Q le nom de *courbe de rosée* (allusion à la rosée du matin). La courbe d'ébullition correspond à l'apparition des premières bulles de vapeur et la courbe de rosée à celle des premières gouttes de liquide. Si le point représentatif se trouve dans la zone de changement d'état, on dit que le *liquide est saturant* ou que la *vapeur est saturante*. Si ce point se trouve sur la courbe d'ébullition, comme la vapeur n'est pas encore apparue, on parle de *liquide saturant seul*. S'il est sur la courbe de rosée, on parle de *vapeur saturante seule*. Enfin, dans le cas où le point représentatif est dans la zone du gaz, la vapeur est dite *vapeur sèche*.

On note ν_l le volume massique du liquide saturant seul et ν_v le volume massique de la vapeur saturante seule. Ils dépendent de la température et correspondent à l'abscisse de l'intersection des courbes de saturation avec l'isotherme considérée. On remarque que ν_v varie beaucoup plus que ν_l entre la température triple T_{Tr} et le point critique et qu'ils sont égaux à T_C . Pour l'eau entre 100 °C et $T_C = 647\text{ °C}$ par exemple :

- à 100 °C : $\nu_l = 1\text{ dm}^3 \cdot \text{kg}^{-1}$ et $\nu_v = 1\,673\text{ dm}^3 \cdot \text{kg}^{-1}$,
- à 647 °C : $\nu_l = \nu_v = 3,1\text{ dm}^3 \cdot \text{kg}^{-1}$,

d'où une variation beaucoup plus importante pour ν_v que pour ν_l . D'autre part, loin du point critique, on pourra en général négliger le volume massique du liquide devant celui du gaz.

2.3 Titres massiques et molaires (PCSI)

On peut facilement déterminer la composition du mélange liquide-vapeur par simple lecture du diagramme de Clapeyron. Pour cela, on définit les titres en liquide et en vapeur. On note n_l et m_l (respectivement n_v et m_v) le nombre de moles et la masse du liquide (respectivement de la vapeur), n le nombre de moles total et m la masse totale.

On définit les titres massiques par :

$$y_l = \frac{m_l}{m_l + m_v} \quad \text{pour le liquide,} \quad y_v = \frac{m_v}{m_l + m_v} \quad \text{pour la vapeur,} \quad (40.1)$$

et les titres molaires par :

$$x_l = \frac{n_l}{n_l + n_v} \quad \text{pour le liquide,} \quad x_v = \frac{n_v}{n_l + n_v} \quad \text{pour la vapeur.} \quad (40.2)$$

Ces grandeurs sont sans dimension et ont pour somme 1 :

$$\gamma_v + \gamma_l = 1 \quad \text{et} \quad x_v + x_l = 1$$

Étant donné que la masse molaire M est la même pour les deux phases, $m_l = Mn_l$ et $M_v = Mn_v$, les titres molaires et massiques sont identiques :

$$\gamma_v = x_v \quad \text{et} \quad \gamma_l = x_l$$

Comment peut-on déterminer les titres massiques d'un mélange représenté par le point M sur le diagramme de Clapeyron (Cf. figure 40.9) ?

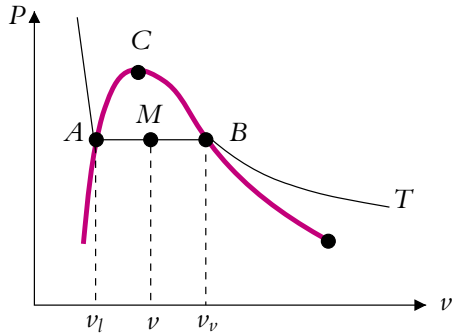


Figure 40.9 Détermination des titres de M .

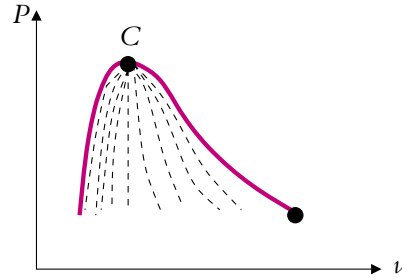


Figure 40.10 Courbes isotitres en pointillés.

Le volume V_l du liquide est égal à $m_l v_l$, celui de la vapeur V_v à $m_v v_v$ et le volume total V , égal à la somme des deux, vaut aussi mv . Pour M , on a donc le système d'équations suivant :

$$\begin{cases} mv = m_v v_v + m_l v_l \\ m = m_v + m_l \end{cases} \quad (40.3)$$

La résolution donne :

$$m_l = m \frac{v_v - v}{v_v - v_l} = m \frac{MB}{AB} \quad (40.4)$$

et

$$m_v = m \frac{v - v_l}{v_v - v_l} = m \frac{AM}{AB} \quad (40.5)$$

Ainsi, par simple lecture de la longueur des segments AB , AM et MB , on peut déterminer : $\gamma_v = \frac{AM}{AB}$ et $\gamma_l = \frac{MB}{AB}$. On peut aussi exprimer le volume massique en M avec les titres :

$$v = \frac{m_l}{m} v_l + \frac{m_v}{m} v_v \Rightarrow v = \gamma_l v_l + \gamma_v v_v = (1 - \gamma_v) v_l + \gamma_v v_v = \gamma_l v_l + (1 - \gamma_l) v_v \quad (40.6)$$

Pour faciliter la détermination des titres, on trace souvent des courbes dites isotitres directement sur le diagramme en utilisant les relations précédentes (Cf. figure 40.10)

Si le diagramme a pour abscisse les volumes molaires, ce sont les titres molaires que l'on peut déterminer directement.

2.4 Expérience des tubes de Natterer

Cette expérience met en évidence le comportement du fluide au voisinage du point critique. Il faut choisir un fluide dont la température critique n'est pas trop élevée comme le fréon. On dispose de trois tubes numérotés de (1) à (3), de même volume constant mais contenant des masses différentes d'un même corps : $m_1 > m_2 > m_3$ soit $v_1 < v_2 < v_3$. La masse m_2 est choisie de manière à ce que le volume massique v_2 soit égal au volume massique critique. Initialement ces tubes sont à la température T_0 comprise entre T_{Tr} et T_C (Cf. figure 40.11). Les points représentatifs figurent sur le diagramme de la figure 40.12. D'après la position des points de départ, il y a plus de liquide que de gaz dans le tube (1), à peu près autant de liquide et de gaz dans le tube (2) et plus de gaz dans le tube (3).

On chauffe doucement les trois tubes qui évoluent à volume massique constant puisque la masse et le volume de chaque tube sont constants. Les évolutions sont représentées par des droites verticales.

- Tube (1) : le point représentatif se rapproche de la courbe d'ébullition. La quantité de gaz diminue donc le ménisque de séparation liquide/gaz monte jusqu'à disparition complète du gaz.
- Tube (2) critique : le point représentatif reste à peu près à même distance des deux courbes de saturation donc le ménisque ne bouge pas. Lorsque la température atteint T_C , le système passe dans l'état fluide, il n'y plus de différence entre liquide et gaz et le ménisque disparaît brusquement.
- Tube (3) : le point représentatif se rapproche de la courbe de rosée. La quantité de liquide diminue donc le ménisque descend jusqu'à disparition complète du liquide.

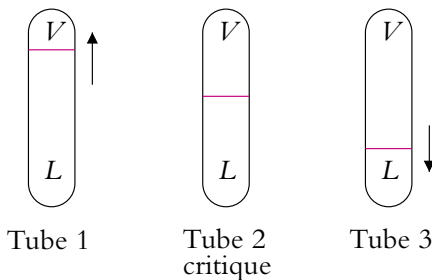


Figure 40.11 Tubes de Natterer.

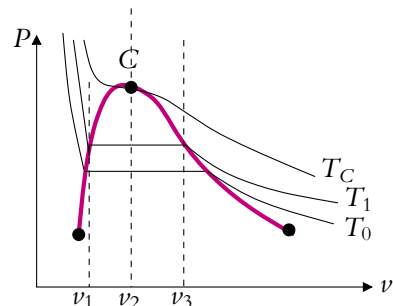


Figure 40.12 Évolution des tubes de Natterer.

En fait, l'expérience est plus spectaculaire lorsqu'ayant réalisé ce qui précède, on laisse les tubes se refroidir à l'air libre, c'est-à-dire lentement. En ce qui concerne les tubes (1) et (3), il y a peu de changement si ce n'est l'évolution inverse des ménisques. Par contre, dans le tube (2) on assiste au phénomène d'*opalescence critique* lorsque le point représentatif passe le point critique. Alors que pour une température supérieure à T_C le corps est translucide et qu'il n'y a pas de ménisque, au passage du point critique le corps devient « opaque » puis, juste après, lorsqu'il redevient translucide, le ménisque réapparaît exactement où il était auparavant. Cette opalescence au passage du point critique est due aux très grandes fluctuations de densité et donc d'indice optique dans le fluide, ce qui diffuse la lumière dans toutes les directions. L'autre phénomène intéressant est l'apparition du ménisque immobile presque au milieu du tube qui prouve, dans ce cas critique, le passage continu de l'état fluide au mélange diphasé.

2.5 Pression de vapeur saturante

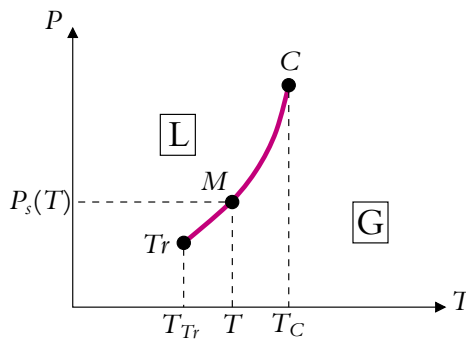


Figure 40.13 Pression de vapeur saturante.

Pour qu'il y ait coexistence de la phase liquide et de la phase vapeur à la température T , il est nécessaire que le point représentatif de l'équilibre soit sur la courbe de changement d'état en M (Cf. figure 40.13). Dans ce cas, la pression ne peut être égale qu'à $P_s(T)$ qui est appelée *pression de vapeur saturante à la température T* . Ainsi, connaissant $P_s(T)$, on peut déterminer, en fonction de la pression P , l'état du corps à la température T .

- Si $P < P_s(T)$, le point M est en-dessous de la courbe et la vapeur est sèche.
- Si $P = P_s(T)$, le point M est sur la courbe et il y a équilibre liquide-vapeur.
- Si $P > P_s(T)$, le point M est au-dessus de la courbe et seul le liquide est stable.

On ne peut donc pas observer de vapeur sèche pour une pression supérieure à la pression de vapeur saturante.

Plus $P_S(T)$ est élevée, plus le corps est volatil. Ainsi l'éther (beaucoup plus volatil que l'eau) a une pression de vapeur saturante égale à 0,59 bar à 293 K alors que celle de l'eau est 0,023 bar. D'après la figure 40.13, P_S est une fonction croissante de la température. Il existe des formules semi-empiriques ou approchées, dont certaines seront établies dans le cours de seconde année (PC/PC* et PT/PT*) :

- Dupré (large domaine de T) $\ln P_S = A - \frac{B}{T} - C \ln T$ (40.7)
- Duperray (autour de 100 °C) $P_S = \left(\frac{t}{100}\right)^4$ avec P en bar et t en °C

On peut interpréter cette variation de P_S en fonction de T en se disant que plus la pression de la vapeur est faible, plus les particules du liquide ayant une vitesse faible peuvent sortir facilement. La vitesse moyenne étant proportionnelle à $\sqrt{T}$, P_S augmente avec la température.

De la même manière, on constate que la température d'ébullition baisse si la pression diminue. Ainsi, en altitude, la température d'ébullition est inférieure à celle au niveau de la mer. Si la pression est 1,013 bar au niveau de la mer, elle n'est plus que de 0,69 bar à 5 000 m et l'eau bout à 88 °C au lieu de 100 °C. Cela signifie qu'en altitude par exemple, il est nécessaire de cuire les pâtes plus longtemps.

Si un fluide se vaporise dans une atmosphère, les expressions sont identiques en remplaçant la pression par la pression partielle¹ du fluide dans l'atmosphère. Pour qu'il y ait équilibre liquide - vapeur, il faut donc que la pression partielle $P_p = xP_{\text{tot}}$ soit égale à la pression de vapeur saturante. À ce propos, on peut se poser la question suivante : pourquoi le linge sèche-t-il quand on le suspend ? Cela signifie que l'eau liquide présente dans le linge se vaporise donc qu'il n'y a pas équilibre entre l'eau liquide dans le linge et la vapeur d'eau de l'air. Si l'air est sec, il contient peu de vapeur d'eau et la pression partielle de l'eau est inférieure à P_S ; seule la vapeur est stable et le liquide se vaporise. Si l'air est très humide (par exemple si il y a du brouillard), l'air est sursaturé en eau, la pression partielle est supérieure à P_S , l'eau du linge ne peut s'évaporer. D'ailleurs, lorsqu'il n'y pas de soleil mais du vent, le linge sèche très bien car le vent évacue la vapeur d'eau et maintient la pression partielle à un niveau bas.

2.6 Exemples

On va traiter deux exemples de détermination de l'état d'un système.

On introduit 4 g d'eau dans un récipient vide de volume $V_0 = 10$ L, indéformable. Dans le premier cas, le récipient est maintenu à 80 °C ($P_S(\text{eau}) = 0,467$ bar). Dans le deuxième cas, la température est 100 °C ($P_S(\text{eau}) = 1,013$ bar).

¹ La pression partielle d'un gaz dans un mélange de gaz est la pression qu'aurait la même quantité de ce gaz s'il occupait seul le volume total du mélange, à la température du mélange.

On considère que la vapeur suit la loi du gaz parfait. Quelle est la composition du mélange à l'équilibre ?

Tout d'abord, on calcule le nombre de mole d'eau sachant que la masse molaire de l'eau est 18 g :

$$n = \frac{4}{18} = 0,22 \text{ mol}$$

a) Premier cas

On fait, par exemple, l'hypothèse que toute l'eau est vaporisée et que la vapeur est sèche. Dans ce cas, si on applique l'équation du gaz parfait, la pression du gaz est :

$$P = \frac{nRT}{V} = \frac{0,22 \times 8,31 \times 353}{10^{-2}} = 0,65 \cdot 10^5 \text{ Pa}$$

La pression de la vapeur sèche serait supérieure à la pression de vapeur saturante donc l'hypothèse est fautive. On est donc amené à faire une seconde hypothèse par exemple un équilibre liquide-vapeur. Dans ce cas, sachant que si toute l'eau était liquide, elle occuperait 4 ml, on peut négliger le volume d'eau liquide devant V . La pression est égale à la pression de vapeur saturante et on peut appliquer l'équation du gaz parfait pour calculer le nombre de moles de gaz :

$$n_v = \frac{P_s V}{RT} = \frac{0,467 \cdot 10^5 \times 10^{-2}}{8,31 \times 353} = 0,16 \text{ mol}$$

Il y a donc 0,16 mole de vapeur et 0,06 mole de liquide. n_v est inférieur à n , ce qui valide l'hypothèse.

► **Remarque :** On évite en général de faire l'hypothèse du liquide seul sauf dans le cas d'une enceinte déformable qui s'adapterait parfaitement au volume occupé par le liquide, sinon il reste toujours un petit volume disponible pour le gaz.

b) Deuxième cas

On fait, cette fois-ci, en premier l'hypothèse qu'il y a équilibre liquide - vapeur. Dans ce cas, le nombre de mole de gaz est :

$$n_v = \frac{P_s V}{RT} = \frac{1,013 \cdot 10^5 \times 10^{-2}}{8,31 \times 373} = 0,32 \text{ mol}$$

On trouve un nombre de moles supérieur au nombre total initial, ce qui est impossible. Il faut donc faire l'hypothèse de la vapeur sèche. Le calcul est similaire à celui du premier cas et on trouve une pression de 0,69 bar qui est cette fois-ci bien inférieure à la pression de vapeur saturante.

3. Les fonctions d'état

3.1 Enthalpie massique de changement d'état

On emploie encore un ancien terme : *chaleur latente de changement d'état* pour *enthalpie massique de changement d'état*.

L'enthalpie massique de vaporisation est la variation d'enthalpie au cours de la vaporisation à pression et à température constante d'un kilogramme de liquide.

On la note l_V ou h_V et elle est positive. En chimie, on utilise l'enthalpie molaire qui correspond à la vaporisation d'une mole.

Si on se réfère à la figure 40.14, l_V correspond à la variation d'enthalpie lors de l'évolution d'un kilogramme de A à B , soit :

$$l_V = \Delta h_{AB} = h_B - h_A \quad (40.8)$$

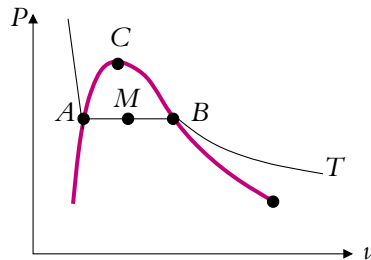


Figure 40.14 Vaporisation à température constante.

S'il y a liquéfaction d'un kilogramme, l'évolution s'effectue de B vers A et la variation d'enthalpie est :

$$\Delta h' = h_A - h_B = -l_V \quad (40.9)$$

Pour une masse m , les variations d'enthalpie sont alors :

- pour une vaporisation $\Delta H = ml_V$,
 - pour une liquéfaction $\Delta H' = -ml_V$.
- (40.10)

S'il y a vaporisation seulement d'une fraction massique γ_v de la masse m , entre A et M par exemple, la variation d'enthalpie sera :

$$\Delta H = H_M - H_A = m\gamma_v l_V \quad (40.11)$$

Pour calculer la variation d'enthalpie, il faut donc faire attention au sens de l'évolution et à la masse qui change d'état.

► Remarques

- On peut définir des enthalpies pour les autres changements d'état : fusion et sublimation, elles aussi positives. La définition correspond toujours à l'évolution de l'état le plus ordonné à l'état le moins ordonné.
- On appelle l_v enthalpie de changement d'état mais il s'agit d'une variation d'enthalpie.
- Comme l'évolution est isobare, $\Delta H = Q$ d'où l'ancien nom de chaleur latente.

3.2 Variation de l'enthalpie de vaporisation l_v avec la température

En fait, l'enthalpie dépend de la pression P et de la température T mais, puisque lors d'un changement d'état, P et T sont liées, l_v ne dépend que de la température T . L'allure de l'enthalpie de vaporisation en fonction de la température est représentée sur la figure 40.15.

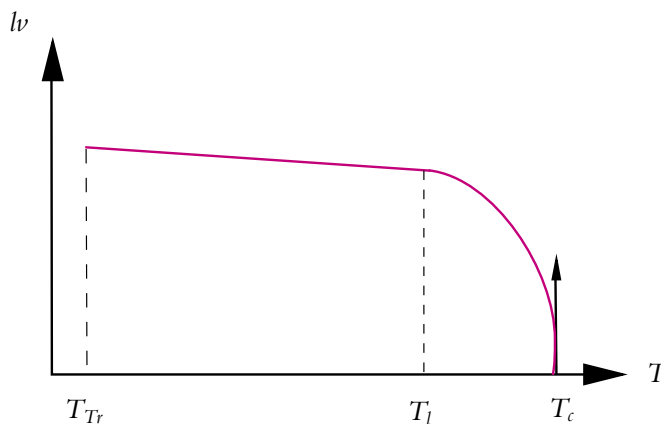


Figure 40.15 Variation de l'enthalpie de vaporisation avec la température.

On remarque une grande zone, jusqu'à la température T_l , dans laquelle la courbe est assimilable à une droite et la variation linéaire. Dans cette zone, on peut appliquer la formule de Regnault :

$$l_v = a - bT \quad (40.12)$$

Pour l'eau, cette formule est applicable jusqu'à $T_l = 200^\circ\text{C}$. Souvent on néglige le terme bT devant a .

Au-delà de la température T_l , on doit ajouter un terme en T^2 à l'expression précédente. Enfin l'enthalpie de vaporisation s'annule pour la température critique T_C et est nulle au-dessus, ce qui confirme le passage continu de la vapeur au liquide quand on contourne le point critique. Il sera montré dans le cours de deuxième année que la pente de l_V est infinie en T_C .

De plus, l'enthalpie de vaporisation n'est définie que pour $T > T_{Tr}$ (Cf. figure 40.15).

3.3 Capacité thermique du liquide saturant seul

On considère le liquide saturant seul, c'est-à-dire un point représentatif sur la courbe d'ébullition. Dans ce cas, on peut appliquer l'hypothèse déjà rencontrée du liquide incompressible et définir c , *capacité thermique massique du liquide saturant seul*, telle que, pour une variation dT de la température le long de la courbe d'ébullition :

$$dU = dH = mc \, dT \quad (40.13)$$

La capacité peut être considérée constante et de O à A (Cf. figure 40.16), on peut écrire :

$$\Delta U = \Delta H = mc(T_A - T_0) \quad (40.14)$$

Quelle est l'erreur commise en confondant ΔU et ΔH ? On calcule les deux variations dans le cas de l'eau entre 0°C et 600°C soit pour 1 kilogramme :

$$\Delta H = \Delta U + \Delta(Pv) \Rightarrow \Delta H = \Delta U + P_S v_l(600^\circ\text{C}) - P_S v_l(0^\circ\text{C}) \quad (40.15)$$

La différence entre les deux est donc due à la variation du produit Pv et, puisque les points sont sur la courbe d'ébullition, il s'agit de la pression de vapeur saturante. On évalue ces termes :

$$P_S v_l(600^\circ\text{C}) = 200.10^5 \times 3.10^{-3} = 6.10^4 \text{ J}$$

$$P_S v_l(0^\circ\text{C}) = 10^5 \times 10^{-3} = 10^2 \text{ J}$$

Soit $\Delta(Pv) \simeq 6.10^4 \text{ J}$. Quant au terme contenant la capacité thermique, il vaut :

$$c(T_{600} - T_0) = 4,18.10^3 \times 600 = 2,4.10^6 \text{ J}$$

En négligeant le terme $\Delta(Pv)$ devant $c\Delta T$, on néglige 6.10^4 devant $2,4.10^6$ et on commet une erreur de 2,5 %, ce qui est acceptable.

3.4 Calcul de l'enthalpie en un point du mélange

Pour calculer l'enthalpie en un point M du mélange, comme cela sera le cas pour l'entropie et l'énergie interne, on part d'un point de référence de la courbe d'ébullition

noté O (Cf. figures 40.16 et 40.17) et on passe par un point de la courbe d'ébullition à la même température que le point M pour lequel on calcule la fonction d'état recherchée.

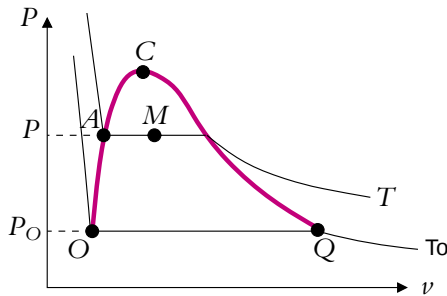


Figure 40.16 Calcul de l'enthalpie d'un mélange liquide - vapeur.

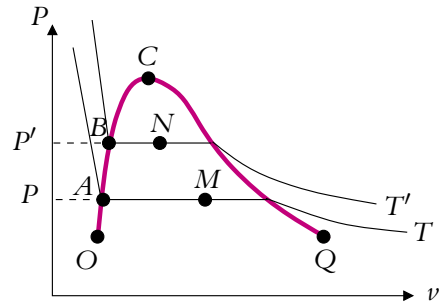


Figure 40.17 Calcul de la variation d'enthalpie au cours d'un changement d'état liquide-vapeur.

Le point O est en général choisi sur l'isotherme triple mais ce n'est pas une obligation. Pour calculer l'enthalpie du système au point M , on décompose en deux étapes, de O à A et de A à M . Le calcul sera fait pour une masse m . Les valeurs des différentes variables dans chaque état sont :

$$\begin{array}{c|c}
 O & \begin{array}{l} T_0 \\ \gamma_l = 1 \\ \gamma_v = 0 \\ P_0 \end{array} \\
 A & \begin{array}{l} T \\ \gamma_l = 1 \\ \gamma_v = 0 \\ P \end{array} \\
 M & \begin{array}{l} T \\ \gamma_l \\ \gamma_v = 1 - \gamma_l \\ P \end{array}
 \end{array}$$

Alors de O à A , il y a un changement de température du liquide saturant seul :

$$\Delta H_1 = H_A - H_O = mc(T - T_0)$$

De A à M , il y a vaporisation d'une fraction γ_v :

$$\Delta H_2 = H_M - H_A = m\gamma_v l_V = m(1 - \gamma_l)l_V$$

Soit en combinant les deux expressions, on obtient l'enthalpie en M :

$$H_M = H_O + mc(T - T_0) + m\gamma_v l_V \quad (40.16)$$

Il faut se rappeler que puisque H est une fonction d'état, la variation ΔH est indépendante du chemin suivi. Par conséquent, si on désire exprimer la variation d'enthalpie entre deux états représentés par les points M et N , soit on écrit la différence entre deux expressions telles que (40.16), soit on raisonne directement en décomposant la transformation comme ce qui suit. De M à A , il y a liquéfaction de la fraction γ_v qui

était sous forme de vapeur :

$$\Delta H_1 = H_A - H_M = -m\gamma_v l_V(T)$$

De A à B , il y a changement de température du liquide saturant seul :

$$\Delta H_2 = H_B - H_A = mc(T' - T)$$

De B à N , il y a vaporisation d'une fraction γ'_v :

$$\Delta H_3 = H_N - H_B = m\gamma'_v l_V(T')$$

Soit :

$$\begin{aligned} \Delta H &= H_N - H_M = \Delta H_1 + \Delta H_2 + \Delta H_3 \\ &= -m\gamma_v l_V(T) + mc(T' - T) + m\gamma'_v l_V(T') \end{aligned} \quad (40.17)$$

Entre deux points voisins d'une même isotherme, de γ_v à $\gamma_v + d\gamma_v$, on obtient :

$$dH = m d\gamma_v l_V \quad (40.18)$$

3.5 Calcul de l'énergie interne en un point du mélange

On procède de la même manière que précédemment pour l'état M de la figure 40.16.

De O à A , il y a changement de température du liquide saturant seul et on confond U et H :

$$\Delta U_1 = U_A - U_O = mc(T - T_0)$$

De A à M , on vaporise une fraction γ_v à la température T , à la pression de vapeur saturante constante $P_S(T) = P$:

$$\Delta U_2 = \Delta H_2 - \Delta(PV) = m\gamma_v l_V - P(V_M - V_A)$$

En A , le liquide est seul donc le volume V_A est égal à $m v_l$. Pour V_M , on utilise l'expression (40.6), soit :

$$V_M = m(\gamma_l v_l + (1 - \gamma_l) v_v)$$

Ainsi on obtient :

$$\Delta U_2 = m\gamma_v l_V + P_S(1 - \gamma_l)(v_l - v_v)$$

L'énergie interne en M est donc :

$$U_M = U_0 + mc(T - T_0) + m(1 - \gamma_l)(l_V(T) + P_S(T)(v_l - v_v)) \quad (40.19)$$

Il arrive souvent qu'on néglige v_l devant v_v dans l'expression précédente.



$Q = \Delta H$ et donc $Q = l_V$ uniquement si le changement d'état se fait à pression constante. Si, par exemple, on met une goutte de liquide dans une enceinte vide indéformable, la vaporisation aura lieu à volume constant. Dans ce cas, $Q = \Delta U$ mais, par définition, on aura toujours $\Delta H = l_V$ et $Q \neq l_V$.

3.6 Calcul de l'entropie en un point du mélange

Comme précédemment, on part d'un état de référence O et on décompose la transformation.

De O à A , il s'agit d'un échauffement du liquide saturant seul :

$$dS_1 = mc \frac{dT}{T} \Rightarrow \Delta S_1 = S_A - S_O = mc \int_{T_0}^T \frac{dT}{T} = mc \ln \left(\frac{T}{T_0} \right)$$

De A à M , la transformation correspond à la vaporisation à P et T constantes d'une fraction γ_v . **Cette transformation est réversible.** On peut s'en convaincre en imaginant un cylindre plongé dans un thermostat fermé d'un côté par un piston. Ce cylindre contient un mélange diphasé pour un certain volume. Si on déplace lentement le piston, tout en maintenant la pression constante, on peut faire varier la composition du mélange à loisir. Puisque la transformation est isobare, le transfert thermique est égal à la variation d'enthalpie. Alors la variation d'entropie vaut :

$$dS_2 = \frac{\delta Q_{\text{rev}}}{T} = \frac{dH}{T} \Rightarrow \Delta S_2 = S_M - S_A = \int_A^M \frac{dH}{T} = \frac{\Delta H}{T} = \frac{m\gamma_v l_V}{T}$$

On en déduit l'entropie en M :

$$S_M = S_O + mc \ln \left(\frac{T}{T_0} \right) + \frac{m\gamma_v l_V}{T} \quad (40.20)$$

Si, comme pour l'enthalpie, on veut calculer la variation d'entropie de M à N pour la figure 40.17, on obtient :

$$\Delta S_{MN} = S_N - S_M = \frac{-m\gamma_v l_V(T)}{T} + mc \ln \left(\frac{T'}{T} \right) + \frac{m\gamma'_v l_V(T')}{T'} \quad (40.21)$$

Sur une isotherme, pour une variation dy_v du titre en vapeur, la variation élémentaire d'entropie vaut dS :

$$dS = \frac{ml_V}{T} dy_v$$

L'expression précédente montre que si le titre en vapeur augmente, c'est-à-dire $dy_v > 0$, alors, $dS > 0$ et l'entropie augmente. En effet, puisque la quantité de vapeur est plus importante, on a moins d'informations sur la position des particules donc l'entropie doit augmenter.

4. Expression des fonctions d'états à partir de valeurs tabulées

Il est très utile de savoir écrire une fonction d'état à partir des données des tables dans lesquelles on trouve soit des grandeurs massiques soit des grandeurs molaires. On note en lettres minuscules les grandeurs massiques avec les indices habituels l (liquide) et v (vapeur). On obtient, comme pour le volume V (voir paragraphe 2.3) :

$$\begin{aligned}U &= m(\gamma_l u_l + \gamma_v u_v) \\H &= m(\gamma_l h_l + \gamma_v h_v) \\S &= m(\gamma_l s_l + \gamma_v s_v)\end{aligned}\tag{40.22}$$

A. Applications directes du cours

1. Correction d'affirmations erronées

Expliquer pourquoi les phrases suivantes sont fausses et, si possible, les rectifier pour les rendre exactes.

1. Un changement d'état a lieu à température et pression constantes.
2. Au cours d'un changement d'état isotherme d'un corps pur, l'énergie interne ne varie pas.
3. Au cours de la vaporisation isotherme d'un kilogramme d'un corps pur, la variation d'énergie interne est égale à l'enthalpie massique de vaporisation.
4. Au point triple, l'état d'un système est parfaitement déterminé.

2. Canette autoréfrigérante

A l'occasion de la Coupe du Monde de Football 2002, une canette autoréfrigérante a été mise au point. Elle comprend un réservoir en acier contenant le liquide réfrigérant. Lorsqu'on ouvre la canette, ce liquide est libéré, il se détend brusquement et se vaporise en traversant une spirale en aluminium qui serpente à travers la boisson à refroidir. Le volume de la boisson à refroidir est 33 mm^3 , et l'on considèrera pour simplifier qu'il s'agit d'eau de capacité thermique massique $c = 4,2 \text{ kJ.kg}^{-1}$. On considère que le corps réfrigérant est constitué d'une masse $m_r = 60 \text{ g}$ de N_2 dont l'enthalpie massique de vaporisation est $L_v = 200 \text{ kJ.kg}^{-1}$. Calculer la variation de température de la boisson.

3. Fonte de glace dans l'eau

Dans un récipient parfaitement calorifugé, on met un morceau de glace à la température de 0°C dans un kilogramme d'eau initialement à la température de 20°C .

1. Déterminer la masse minimale de glace nécessaire pour que l'eau soit à la température de 0°C .
2. Calculer ΔS_e la variation d'entropie de l'eau initialement à l'état liquide.
3. Même question pour ΔS_{gl} pour l'eau initialement sous forme de glace.
4. En déduire le bilan d'entropie de l'évolution. Conclure.

On donne la capacité thermique massique de l'eau $c = 4,2 \cdot 10^3 \text{ J.K}^{-1}.\text{kg}^{-1}$ et la chaleur latente massique de fusion de la glace $L = 336 \cdot 10^3 \text{ J.kg}^{-1}$.

4. Étude de quelques transformations d'un corps, d'après CCP TSI 2006

Dans cet exercice, on s'intéresse à l'eau dont le diagramme des phases est donné par :

1. Reproduire ce diagramme et le compléter en donnant les domaines d'existence des différentes phases et en définissant les points caractéristiques.
2. Définir la pression de vapeur saturante et préciser de quel(s) paramètre(s) elle dépend.

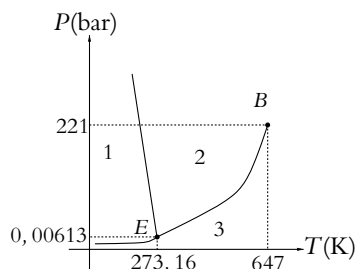


Figure 40.18

3. Comment appelle-t-on le passage de la vapeur au liquide ?
4. Représenter le diagramme donnant la pression en fonction du volume pour la transformation correspondante. On définira les domaines et on tracera les courbes de rosée et d'ébullition.
5. Soit une enceinte cylindrique diathermane de volume initial V , ce volume pouvant être modifié en déplaçant sans frottement un piston. L'ensemble est maintenu sous la pression atmosphérique à la température $T = 373 \text{ K}$. A cette température la pression de vapeur saturante vaut 1,0 bar. La vapeur d'eau sèche et saturante sera considérée comme un gaz parfait. On néglige le volume occupé par la phase liquide devant le volume occupé par la vapeur, ainsi le volume de la phase gazeuse est égal au volume total de l'enceinte. Le cylindre étant initialement vide, on introduit, piston bloqué, une masse m d'eau. Déterminer la masse maximale $m_{\max}$ d'eau qu'on peut introduire pour que l'eau soit entièrement sous forme vapeur. On donnera sa valeur en fonction de R , T , V , P_S la pression de vapeur saturante et M_{eau} la masse molaire de l'eau.
6. On considère que la masse d'eau introduite est inférieure à $m_{\max}$. Dans quel état se trouve l'eau ?
7. Pour obtenir l'équilibre entre les phases liquide et vapeur de l'eau, faut-il augmenter ou diminuer le volume ? Déterminer le volume limite $V_{\lim}$ à partir duquel on a cet équilibre.
8. La masse m d'eau introduite est telle qu'on a l'équilibre entre les phases liquide et vapeur. Déterminer la fraction massique d'eau sous forme vapeur.

5. Vaporisation de l'azote

On place de l'azote liquide dans un récipient posé sur une balance. Il est possible de chauffer le liquide par une résistance parcourue par un courant de 2,95 A et soumise à une tension de 16,5 V. On mesure alors l'évolution de la masse de l'azote au cours du temps :

temps (min)	masse (g)	temps (min)	masse (g)
0	415	9	361
1	411	10	345
2	409	11	329
3	406	12	314
4	403	13	298
5	400	14	296
6	397	15	294
7	394	16	292
8	377	17	289

1. Tracer la masse d'azote en fonction du temps.
2. Analyser cette évolution.

3. Montrer que les données précédentes permettent de déterminer l'enthalpie de vaporisation de l'azote.
4. Obtenir sa valeur numérique à partir des données.

B. Exercices et problèmes

1. Formation de la neige artificielle, d'après CCP PC 2001.

La neige artificielle est obtenue en pulvérisant de fines gouttes d'eau liquide à $T_1 = 10^\circ\text{C}$ dans l'air ambiant à $T_a = -15^\circ\text{C}$.

1. Dans un premier temps la goutte d'eau supposée sphérique (rayon $R = 0,2\text{ mm}$) se refroidit en restant liquide. Elle reçoit de l'air extérieur un transfert thermique $h(T_a - T(t))$ par unité de temps et de surface, où $T(t)$ est la température de la goutte. On rappelle que la masse volumique de l'eau est $\rho = 10^3\text{ kg.m}^{-3}$ et sa capacité thermique massique $c = 4,18.10^3\text{ J.kg}^{-1}$.

- a) Établir l'équation différentielle vérifiée par $T(t)$.
- b) En déduire le temps t auquel $T(t)$ est égale à -5°C . Effectuer l'application numérique pour $h = 65\text{ W.m}^{-2}.\text{K}^{-1}$.

2. a) Lorsque la goutte atteint la température de -5°C , la surfusion cesse : la goutte est partiellement solidifiée et la température devient égale à 0°C . Calculer la fraction x de liquide restant à solidifier en supposant la transformation très rapide et adiabatique. On néglige aussi la variation de volume. L'enthalpie massique de fusion de la glace est $L_f = 335.10^3\text{ J.kg}^{-1}$.

- b) Au bout de combien de temps la goutte est-elle complètement solidifiée ?

2. Vaporisation réversible ou irréversible

On vaporise une masse $m = 1\text{ g}$ d'eau liquide des deux manières suivantes :

- la masse m est enfermée à 100°C sous la pression atmosphérique, dans un cylindre fermé par un piston. Par déplacement lent du piston, on augmente le volume à température constante et on s'arrête dès que toute l'eau est vaporisée. Le volume est alors égal à $V = 1,67\text{ L}$.
- On introduit rapidement la masse m d'eau liquide initialement à 100°C dans un récipient fermé de même température, de volume $1,67\text{ L}$, initialement vide. L'enthalpie massique de vaporisation de l'eau est $L_v = 2,25.10^6\text{ J.kg}^{-1}$.

1. Pour chacun des processus précédents, calculer le transfert thermique fourni par le thermostat et les variations d'énergie interne, d'enthalpie et d'entropie de l'eau.
2. Calculer l'entropie créée lors du processus irréversible.

3. Machine frigorifique

On considère une machine frigorifique fonctionnant suivant le cycle de transformations quasi-statiques d'un fluide $ABCD$ où AB et CD sont des isentropiques, BC une isotherme à $T_1 = -5^\circ\text{C}$ et DA une isotherme à $T_2 = 15^\circ\text{C}$ (D et A sont aux extrémités du palier de liquéfaction).

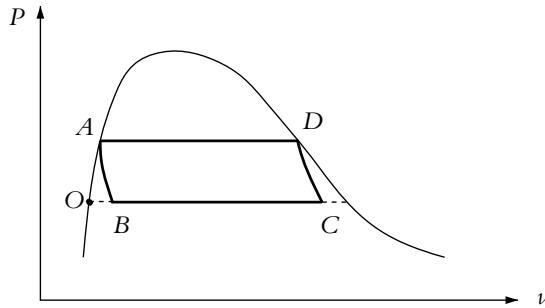


Figure 40.19

On note $c_L = 4,18 \text{ kJ.kg}^{-1}$ la capacité thermique massique du fluide le long de la courbe d'ébullition, $L_v(T_1)$ et $L_v(T_2)$ les enthalpies de changements d'état à T_1 et T_2 . On donne $L_v(T_2) = 1,30 \cdot 10^6 \text{ J.kg}^{-1}$

1. La masse totale du système étant $m = 1 \text{ kg}$ de fluide, exprimer la masse de vapeur en B et en C.
2. Calculer le travail W reçu au cours du cycle.
3. Calculer l'efficacité du réfrigérateur fonctionnant suivant ce cycle, définie par $\eta = \left| \frac{Q_{BC}}{W} \right|$.

4. Étude de l'air humide, d'après CCP PC 2003

L'air humide est un mélange d'air sec et de vapeur d'eau, les caractéristiques de l'air humide étant liées aux proportions de chacun des constituants. Dans ce problème, on étudie l'air humide sous la pression atmosphérique $P = 1,013 \cdot 10^5 \text{ Pa}$ en supposant que l'air sec et la vapeur d'eau se comportent comme des gaz parfaits. On donne la masse molaire de l'air sec $M_a = 29 \text{ g.mol}^{-1}$, celle de l'eau $M_v = 18 \text{ g.mol}^{-1}$, la capacité thermique massique de l'air sec $c_{pa} = 1006 \text{ J.K}^{-1}.\text{kg}^{-1}$, celle de l'eau sous forme de gaz $c_{pv} = 1923 \text{ J.K}^{-1}.\text{kg}^{-1}$, la chaleur latente massique de vaporisation de l'eau $l = 2500 \text{ kJ.kg}^{-1}$ et les pressions de vapeur saturante de l'eau en fonction de la température :

$\theta(^{\circ}\text{C})$	0,0	5,0	10	15	20	25	30	40	45
$P_{vs}(\text{Pa})$	610	880	1227	1706	2337	3173	4247	7477	9715

On définit la pression partielle P_a de l'air sec dans un volume V d'air humide à la température T comme la pression de l'air sec dans le même volume et à la même température. On fait de même pour la pression partielle de la vapeur d'eau P_v . Compte tenu de l'hypothèse que les deux sont assimilés à des gaz parfaits, la pression P de l'air humide vérifie $P = P_a + P_v$.

1. Soit m_a la masse d'air sec contenue dans le volume V d'air humide à la température T . Montrer qu'on peut écrire $P_a V = m_a R_a T$ où R_a est une constante à exprimer en fonction de R et M_a . Donner sa valeur numérique.
2. Montrer qu'il en est de même pour la vapeur d'eau.

3. On définit l'humidité spécifique ω de l'air humide à la température T comme le rapport de la masse de vapeur d'eau contenue dans le volume V d'air humide à la masse d'air sec contenue dans ce même volume. Établir qu'elle s'exprime sous la forme $\omega = A \frac{P_v}{P - P_v}$ en donnant l'expression et la valeur numérique de la constante A .
 4. La sensation de se trouver dans un air sec ou humide est liée au degré hygrométrique ou humidité relative $\epsilon = \frac{P_v}{P_{vs}}$. Calculer les masses m_a d'air sec et m_v de vapeur d'eau d'un mètre cube d'air humide à 15°C et de degré hygrométrique 0,85.
 5. Un air humide tel que $\epsilon = 1$ ne peut plus accepter d'eau sous forme vapeur. L'eau supplémentaire se présente alors sous forme de gouttelettes d'eau suffisamment fines pour rester en suspension formant ainsi un brouillard. Tracer précisément la courbe représentative de l'air humide saturé dans le diagramme de Carrier donnant ω en fonction la température θ . Indiquer, en la justifiant, la zone de brouillard sur ce diagramme.
 6. La température de rosée est la température de l'air humide saturé en humidité. On peut la mesurer par un hygromètre à condensation : on place dans l'air humide une petite surface dont on fait varier la température jusqu'à l'apparition sur celle-ci de condensat (rosée ou buée), la température de la surface est alors celle du point de rosée. Calculer le degré hygrométrique d'un air humide à une température de 30°C dont la température de rosée est égale à 10°C .
 7. L'enthalpie de l'air humide tient compte de l'enthalpie de ses constituants définie sur la base des conventions suivantes : l'origine de l'enthalpie de l'air sec est prise à 0°C et l'enthalpie de référence pour l'eau est celle de l'eau liquide à 0°C . Donner l'expression de l'enthalpie massique h de l'air humide en fonction de m_a , m_v , c_{pa} , c_{pv} , l et θ la température en degrés Celsius de l'air humide.
 8. Même question pour H^* l'enthalpie spécifique de l'air humide contenant un kilogramme d'air sec en fonction de ω , c_{pa} , c_{pv} , l et θ .
 9. Les techniques de climatisation et de conditionnement de l'air reposent sur des opérations de mélange, d'échauffement, de refroidissement ou d'humidification de l'air humide dans le but d'améliorer les conditions de confort. Les mélanges d'airs humides sont également à l'origine de phénomènes météorologiques comme les brouillards côtiers issus de la rencontre de l'air frais et sec provenant de l'intérieur des terres et de l'air marin fortement humide ou encore la formation des nuages issue de l'ascension de l'air humide plus léger que l'air sec et de sa rencontre avec l'air plus froid en altitude.
- Un réchauffeur fournit un transfert thermique Q à pression constante à un kilogramme d'air humide initialement à la température θ , d'humidité spécifique ω et contenant une masse m_a d'air sec. Exprimer l'humidité spécifique ω_q ainsi que la température θ_q de l'air humide obtenu.
10. Décrire qualitativement l'évolution du degré hygrométrique lors de l'opération de chauffage.
 11. Une installation de climatisation industrielle assure le réchauffage isobare d'un débit massique de $12 \cdot 10^3 \text{ kg} \cdot \text{h}^{-1}$ d'un air humide entrant à une température de 15°C avec un degré hygrométrique $\epsilon = 0,85$. La température de sortie est de 45°C . Quelle est la puissance du réchauffeur ?
 12. On mélange deux airs humides de température θ_1 et θ_2 , d'humidité spécifique ω_1 et ω_2 contenant respectivement les masses d'air sec m_{a1} et m_{a2} . Établir les relations permettant de

calculer en fonction de ω_1 , ω_2 , m_{a1} , m_{a2} , H_1^* et H_2^* l'humidité spécifique ω_3 et l'enthalpie spécifique H_3^* du mélange.

13. Application numérique : mélange de un kilogramme d'air humide dans les états 1 tel que $\theta_1 = 35^\circ\text{C}$ et $\omega_1 = 0,035$ et 2 tel que $\theta_2 = 25^\circ\text{C}$ et $\omega_2 = 0,0039$. Donner les valeurs numériques de ω_3 et θ_3 .

14. Même question pour les états 4 tel que $\theta_4 = 40^\circ\text{C}$ et $\epsilon_4 = 1$ et 5 tel que $\theta_5 = 5^\circ\text{C}$ et $\epsilon_5 = 0,2$. On mélange des quantités d'air humide contenant chacune un kilogramme d'air sec. Où se trouve le point obtenu dans le diagramme de Carrier ? Conclure.

15. Un thermomètre indique une température supérieure de 3°C à celle obtenue à la question précédente. Quelle est la raison de cet écart ? A l'aide du diagramme de Carrier, en déduire la masse m_e d'eau liquide présente dans le mélange.

Le fonctionnement de nombreux appareils thermiques peut être décrit par la thermodynamique : par exemple, les moteurs à essence et diesel, les réfrigérateurs, les pompes à chaleur, les centrales électriques, les usines d'incinération... Ces machines thermiques font partie de la vie courante, ce qui explique l'importance de leur étude.

1. Description

Une *machine thermique* est constituée d'un système (M) qui, en général, décrit des cycles fermés successifs. Ce système échange du transfert thermique avec les thermostats (notée Σ_i) et du travail avec les systèmes mécaniques (notés SM_i) en contact avec lui. On modélise donc l'ensemble par le schéma de la figure 41.1, avec la convention d'orientation habituelles (Q et W reçus par (M))

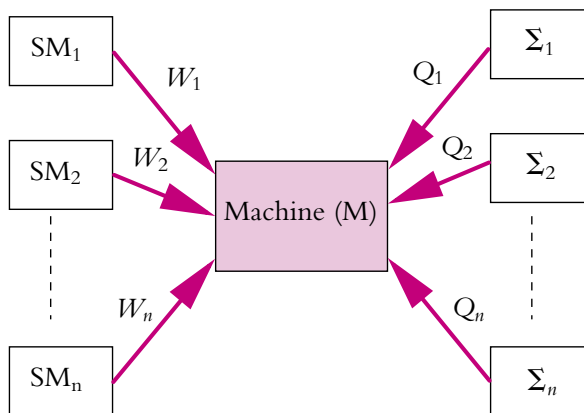


Figure 41.1 Schéma général d'une machine thermique.

Les systèmes mécaniques (ils peuvent être aussi électriques) sont supposés parfaits et n'échangent que du travail avec le système (M). Les sources de chaleur sont soit parfaites (elles fonctionnent alors de manière réversibles et isothermes) soit non parfaites c'est-à-dire qu'elle ont une capacité thermique finie mais sont supposées fonctionner de manière isotherme sur un cycle.

2. Machine monotherme

2.1 Travail maximal récupérable

On étudie le fonctionnement d'une machine monotherme, c'est-à-dire en contact avec un seul thermostat supposé ici parfait (Cf. figure 41.2).

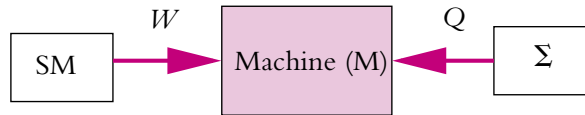


Figure 41.2 Schéma synoptique d'une machine monotherme.

Quel travail maximal peut fournir le système (M) ?

Le premier principe appliqué à (M) donne :

$$\Delta U_M = W + Q$$

Le système (M) n'échange du transfert thermique qu'avec (Σ) donc l'entropie échangée par (M) est :

$$S_{\text{éch}} = \frac{Q}{T_\Sigma}$$

L'entropie créée s'écrit alors :

$$S_{\text{créée}} = \Delta S_M - \frac{Q}{T_\Sigma} \geq 0$$

En combinant les précédentes expressions pour éliminer Q et exprimer W , on obtient :

$$\Delta S_M - \frac{\Delta U_M - W}{T_\Sigma} \geq 0 \Leftrightarrow W \geq \Delta U_M - T_\Sigma \Delta S_M \quad (41.1)$$

Dans le cas d'une transformation réversible, l'inégalité (41.1) devient une égalité.

Le travail fourni à l'extérieur W' est l'opposé de W soit d'après la relation (41.1) :

$$W' \leq -(\Delta U_M - T_\Sigma \Delta S_M) \Rightarrow W'_{\text{max}} = -\Delta U_M + T_\Sigma \Delta S_M \quad (41.2)$$

Le travail maximal récupérable $W'_{\max}$ est donné par l'équation précédente et il est atteint dans le cas réversible.

2.2 Principe de Carnot (1824)

On considère une machine monotherme décrivant des cycles. Dans ce cas, l'expression (41.2) s'applique sur un nombre entier de cycles : $\Delta U_M = 0$ et $\Delta S_M = 0$. On en déduit :

$$W \geq 0 \quad (41.3)$$

Une machine monotherme ne peut que recevoir du travail. Ainsi le moteur monotherme n'existe pas.

Comme cela a déjà été évoqué dans le chapitre sur le deuxième principe, le premier principe n'est pas suffisant pour prouver l'inexistence du moteur monotherme.

La version historique du second principe énoncé par Carnot, et portant le nom de *principe de Carnot*, est la suivante :

Pour qu'un système décrivant un cycle fournisse du travail, il doit nécessairement échanger du transfert thermique avec au moins deux sources à des températures différentes.

Il est aussi rapide quand on n'a pas démontré l'expression générale (41.1) de refaire la démonstration directement avec $\Delta U_M = 0$ et $\Delta S_M = 0$.

3. Étude d'une machine ditherme

La machine ditherme est une machine en contact avec deux thermostats. Bien que la machine ditherme ne soit pas la seule rencontrée en pratique, il s'agit de la plus courante et on peut s'y ramener. C'est pour cela qu'on l'étudie en particulier. C'est la machine la plus simple pouvant fonctionner en moteur ou récepteur puisque, d'après le principe de Carnot, la machine monotherme ne peut fonctionner qu'en récepteur.

3.1 Inégalité de Clausius (vers 1850)

Cette inégalité a déjà été démontrée dans le cas général. On reprend la démonstration dans le cas d'une machine ditherme (Cf. figure 41.3).

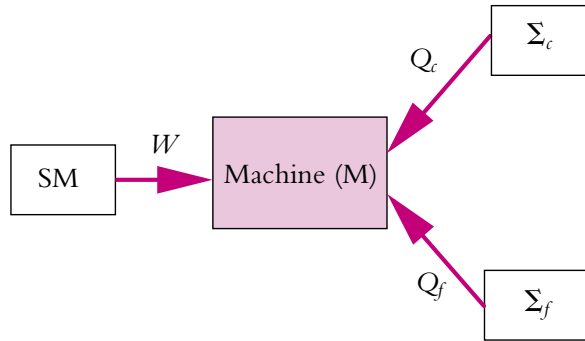


Figure 41.3 Schéma synoptique d'une machine ditherme.

On fait l'hypothèse ici que la machine (M) décrit des cycles donc sur un nombre entier de cycles, $\Delta U_M = 0$ et $\Delta S_M = 0$. Si on applique le deuxième principe à (M) :

$$\Delta S_M = S_{\text{éch}} + S_{\text{créée}} = 0 \Rightarrow S_{\text{éch}} = -S_{\text{créée}}$$

Or puisque (M) n'est en contact qu'avec les deux thermostats Σ_c et Σ_f parfaits, dont les températures T_{Σ_c} et T_{Σ_f} sont constantes, l'entropie échangée s'écrit :

$$S_{\text{éch}} = \frac{Q_c}{T_{\Sigma_c}} + \frac{Q_f}{T_{\Sigma_f}}$$

Ainsi :

$$S_{\text{créée}} \geq 0 \Rightarrow S_{\text{éch}} \leq 0$$

d'où :

$$\frac{Q_c}{T_{\Sigma_c}} + \frac{Q_f}{T_{\Sigma_f}} \leq 0 \quad (41.4)$$

Cette inégalité porte le nom d'inégalité de Clausius.

► Remarques

- Il ne faut pas oublier que l'inégalité de Clausius se démontre à partir du deuxième principe.
- Dans le cas de cycles réversibles, l'inégalité de Clausius devient une égalité.
- Dans le cas de n sources, par le même raisonnement, on trouve : $\sum_{k=0}^n \frac{Q_k}{T_{\Sigma_k}} \leq 0$.

Attention, l'inégalité (41.4) **n'est valable que pour des thermostats parfaits, c'est-à-dire évoluant à température constante**. Dans le cas de thermostats de capacité C non infinie, il faut revenir au deuxième principe et exprimer correctement l'entropie échangée par (M) ou, ce qui revient au même, la variation d'entropie du

thermostat. Par exemple, si le thermostat (Σ_c) a une capacité thermique C_{Σ_c} et que sa température passe de T_c à T'_c , l'entropie échangée par (M) avec (Σ_c) sera :

$$S_{\text{éch},1} = \int \frac{\delta Q_c}{T_{\Sigma_c}} = \int \frac{-C_{\Sigma_c} dT_{\Sigma_c}}{T_{\Sigma_c}} = -C_{\Sigma_c} \ln \left(\frac{T'_c}{T_c} \right)$$

Attention au signe $-$: C_{Σ_c} est la capacité thermique du thermostat donc $C_{\Sigma_c} dT_{\Sigma_c}$ est le transfert thermique **reçu par le thermostat**.

Globalement, l'entropie échangée est donc :

$$S_{\text{éch}} = -C_{\Sigma_c} \ln \left(\frac{T'_c}{T_1} \right) + \frac{Q_f}{T_{\Sigma_f}}$$

et puisque comme précédemment, $S_{\text{éch}} \leq 0$, on obtient :

$$-C_{\Sigma_c} \ln \left(\frac{T'_c}{T_c} \right) + \frac{Q_f}{T_{\Sigma_f}} \leq 0 \quad (41.5)$$

Encore une fois, il ne faut pas oublier que l'inégalité de Clausius provient du deuxième principe. Historiquement, elle a même constitué un énoncé du deuxième principe.

3.2 Diagramme de Raveau

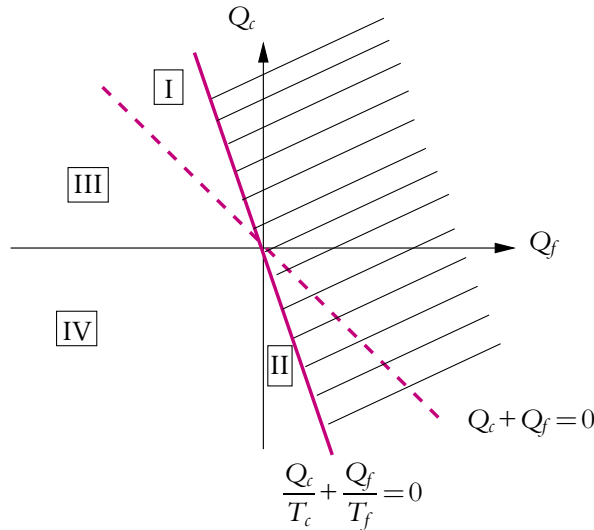


Figure 41.4 Diagramme de Raveau.

Il s'agit d'une représentation en deux dimensions des zones utilisables d'une machine ditherme. On porte en abscisse le transfert thermique reçu par CM de la part de source froide Q_f et en ordonnée celui reçu de la part de source chaude Q_c (Cf. figure 41.4).

Le deuxième principe implique, par l'intermédiaire de l'inégalité de Clausius, que $\frac{Q_c}{T_c} + \frac{Q_f}{T_f} \leq 0$ donc le demi-plan au-dessus de la droite $Q_c = -\frac{T_c}{T_f} Q_f$ est une zone interdite (zone hachurée sur le diagramme).

Il reste à déterminer quelles sont les zones intéressantes dans le demi-plan inférieur. Si on s'intéresse à une machine fonctionnant suivant des cycles, le premier principe donne : $\Delta U = 0 \Rightarrow Q_c + Q_f + W = 0 \Rightarrow W = -(Q_c + Q_f)$. La droite d'équation $Q_c + Q_f = 0$ va donc délimiter les zones de comportement moteur (au-dessus de la droite, $W < 0$) et récepteur (au-dessous de la droite, $W > 0$). Puisque $\frac{T_c}{T_f} > 1$, la droite résultant de l'inégalité de Clausius a une pente inférieure à l'autre.

Il ne faut pas oublier que les échanges sont orientés vers la machine.

On étudie les quatre zones :

- **Zone I :** $W < 0$, $Q_c > 0$, $Q_f < 0$. Le système se comporte en moteur, il reçoit du transfert thermique de la source chaude et en cède à la source froide. Ce qu'il reçoit en plus par rapport à ce qu'il cède est transformé en travail. Pour un moteur de voiture par exemple, la source froide est l'air et l'échange se fait par l'intermédiaire du radiateur. La source chaude n'est pas à température constante, l'énergie provient de l'inflammation du mélange air-carburant.
- **Zone II :** $W > 0$, $Q_c < 0$, $Q_f > 0$. Le système se comporte en récepteur, il reçoit du travail, ce qui lui permet d'effectuer un transfert contraire à l'échange spontané. Il reçoit du transfert thermique de la source froide et en cède à la source chaude. C'est le principe des machines frigorifiques et des pompes à chaleur.
- **Zone III :** $W > 0$, $Q_c > 0$, $Q_f < 0$. C'est une zone non intéressante, car il faut fournir du travail pour effectuer un échange entre source chaude (elle cède de l'énergie) et froide (elle reçoit de l'énergie) qui se fait spontanément sans machine.
- **Zone IV :** $W > 0$, $Q_c < 0$, $Q_f < 0$. C'est encore une zone non intéressante, puisqu'il s'agit de fournir du travail pour apporter du transfert thermique aux deux sources ; autant utiliser une résistance chauffante, au lieu d'une machine thermique sophistiquée.

On étudie dans la suite les rendements des machines thermiques fonctionnant dans les zones I ou II.

4. Exemples de machines dithermes

4.1 Moteur thermique

a) Rendement

On fait l'hypothèse que seule la source chaude est onéreuse.

La définition générale d'un rendement est :

$$\rho = \left| \frac{\text{ce qui est intéressant}}{\text{ce qui coûte}} \right| \quad (41.6)$$

Pour un moteur, c'est le travail qui est intéressant et c'est la source chaude qui est onéreuse. Sachant de plus que $W < 0$ et $Q_c > 0$, le rendement est :

$$\rho = \left| \frac{W}{Q_c} \right| = -\frac{W}{Q_c} \quad (41.7)$$

D'après le premier principe, $W = -(Q_c + Q_f)$, l'expression précédente devient donc :

$$\rho = 1 + \frac{Q_f}{Q_c} \quad (41.8)$$

Il est plus pratique de disposer d'une expression avec des données connues comme les températures. Pour cela, on utilise l'inégalité de Clausius pour majorer $\frac{Q_f}{Q_c}$. Sachant que $Q_c > 0$ et $Q_f < 0$, il vient :

$$\frac{Q_c}{T_c} + \frac{Q_f}{T_f} \leq 0 \Leftrightarrow -\frac{Q_c}{T_c} \geq \frac{Q_f}{T_f} \Leftrightarrow \frac{Q_f}{Q_c} \leq -\frac{T_f}{T_c}$$

D'où en reportant dans (41.8) :

$$\rho = 1 + \frac{Q_f}{Q_c} \leq 1 - \frac{T_f}{T_c} \quad (41.9)$$

Ainsi, le rendement de n'importe quel moteur thermique fonctionnant entre deux sources est inférieur à $1 - \frac{T_f}{T_c}$, expression déjà rencontrée dans le chapitre sur les applications du premier principe et le cycle de Carnot (37.13). Il s'agit du rendement d'un moteur fonctionnant suivant un cycle de Carnot réversible, qu'on note ρ_C , et qui, d'après l'expression (41.9), est atteint par le moteur étudié dans le cas réversible. On peut énoncer le théorème de Carnot.

b) Théorème de Carnot

Tous les moteurs thermiques réversibles fonctionnant entre deux sources à des températures données ont le même rendement $\rho_c = 1 - \frac{T_f}{T_c}$ appelé *rendement de Carnot*. Les machines irréversibles fonctionnant entre ces mêmes sources ont un rendement inférieur à celui des machines réversibles.

On a déjà évoqué l'intérêt de l'étude des cycles réversibles, cet intérêt apparaît véritablement ici. On voit d'autre part que pour avoir le rendement le plus élevé possible, il faut des températures les plus éloignées possibles.

On peut aussi exprimer le rendement en fonction de l'entropie créée. Dans l'étude générale d'une machine ditherme, on a établi que :

$$S_{\text{créée}} = -S_{\text{éch}} = -\left(\frac{Q_c}{T_c} + \frac{Q_f}{T_f} \right)$$

donc :

$$\frac{Q_f}{Q_c} = -\frac{T_f}{T_c} - \frac{S_{créée} T_f}{Q_c}$$

On en déduit le rendement :

$$\rho = 1 - \frac{T_f}{T_c} - \frac{S_{créée} T_f}{Q_c}$$

le dernier terme dans l'expression précédente représente la différence entre le rendement de la machine et le rendement réversible de Carnot.

4.2 Machine frigorifique et pompe à chaleur

Dans les deux cas, le point de fonctionnement est dans la zone II du diagramme de Raveau. Il s'agit de fournir du travail à une machine pour effectuer un transfert thermique dans le sens opposé au sens spontané, c'est-à-dire de la source froide vers la source chaude. Dans le cas de ces deux types de machines thermiques, on ne parle pas de rendement mais d'efficacité car le résultat trouvé est supérieur à 1. On en interprétera la signification à la fin du présent paragraphe.

a) Machine frigorifique

Dans le cas d'une machine frigorifique, la source froide est l'intérieur contenant les denrées à maintenir au froid et la source chaude est la pièce contenant la machine. La grandeur intéressante est le transfert thermique provenant de la source froide et reçu par le fluide circulant dans les canalisations de la machine. Ce qui coûte est le travail électrique qui fait fonctionner la pompe du réfrigérateur. On définit donc l'efficacité par :

$$e = \left| \frac{Q_f}{W} \right| \quad \text{et puisque } Q_f > 0 \text{ et } W > 0, \quad e = \frac{Q_f}{W} \quad (41.10)$$

Soit avec le premier principe, $W = -(Q_c + Q_f)$:

$$e = -\frac{Q_f}{Q_c + Q_f} = -\frac{1}{1 + \frac{Q_c}{Q_f}} \quad (41.11)$$

Comme dans le cas du moteur, on utilise l'inégalité de Clausius pour exprimer l'efficacité en fonction des températures. Avant tout calcul avec des inégalités, il faut prendre garde aux signes des transferts thermiques (ici $Q_c < 0$ et $Q_f > 0$) :

$$\frac{Q_c}{T_c} + \frac{Q_f}{T_f} \leq 0 \quad \Leftrightarrow \quad \frac{Q_c}{Q_f} \leq -\frac{T_c}{T_f} \quad \Leftrightarrow \quad -\left(1 + \frac{Q_c}{Q_f}\right) \geq \frac{T_c}{T_f} - 1 (> 0)$$

D'où l'efficacité :

$$e \leq \frac{1}{\frac{T_c}{T_f} - 1} = \frac{T_f}{T_c - T_f} \quad (41.12)$$

b) Théorème de Carnot

On trouve un résultat semblable à celui des moteurs, à savoir que l'efficacité d'une machine frigorifique ditherme est inférieure à l'efficacité e_c d'une machine frigorifique réversible fonctionnant suivant un cycle de Carnot. Il y a égalité dans le cas réversible.

Contrairement au moteur, il faut que les températures soient les plus proches possible pour avoir l'efficacité la plus grande possible. Cela n'étonne pas, on imagine bien que plus la pièce dans laquelle est le réfrigérateur est froide, plus il sera facile de refroidir l'intérieur.

c) Structure d'un système à condensation

Les systèmes frigorifiques sont en général des systèmes dits *systèmes à condensation* dont le schéma est présenté sur la figure 41.5.

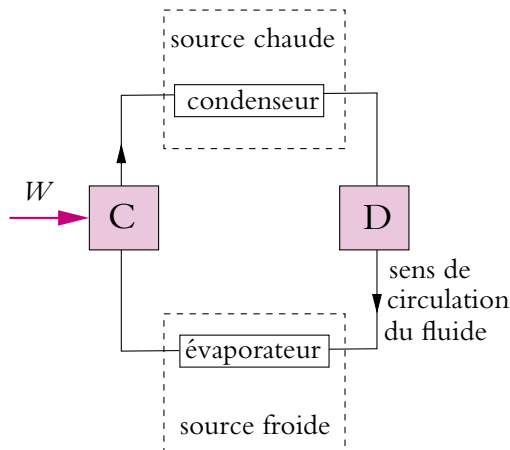


Figure 41.5 Schéma d'un système à condensation.

Le système qu'on appelle « machine » est constitué d'un fluide dit *fluide caloporteur* car c'est lui qui assure les transports d'énergie. Ce fluide circule dans les canalisations que l'on voit à l'intérieur du bac à glaçons et à l'extérieur derrière le réfrigérateur. On utilisait jusqu'à il y a quelques années comme fluides caloporteurs des CFC (chlorofluorocarbène) qui sont maintenant interdits en raison de leur interaction possible avec la couche d'ozone de l'atmosphère puis on a ensuite utilisé des alcanes qui étaient dangereux en cas d'incendie. Actuellement, on emploie certains HCFC (hydrogénéofluorocarbène) qui ont des propriétés proches des CFC mais sont moins stables et sont détruits avant d'arriver dans la zone où se trouve l'ozone.

La source froide est formée de l'air et des denrées à l'intérieur du réfrigérateur, la source chaude est l'air extérieur. Le fluide est en contact avec les deux sources grâce aux tubulures. Il subit une évaporation dans l'évaporateur quand il est en contact

avec la source froide ; cette évaporation nécessite de l'énergie prise à la source froide ($Q_f > 0$). Il subit aussi une liquéfaction dans le condenseur lorsqu'il est en contact avec la source chaude ; il cède alors de l'énergie à la source chaude ($Q_c < 0$).

Le compresseur (noté C sur le schéma) reçoit de l'énergie électrique, il joue le rôle de pompe de manière à faire circuler le fluide mais aussi à le comprimer. Pour se convaincre de la nécessité de la compression, il suffit de revoir au chapitre précédent le diagramme de Clapeyron, sachant que l'isotherme pour le changement d'état à T_c est à une pression plus élevée que celle à T_f . Il faut donc comprimer le fluide quand il passe de la source froide à la source chaude et le détendre (rôle du détendeur noté D sur le schéma) lorsqu'il passe de la source chaude à la source froide.

On peut alors se poser la question suivante, pour se rendre compte si on a bien compris : si on laisse ouverte la porte du réfrigérateur, arrivera-t-on à refroidir la pièce entière ? La réponse est non. La pièce va se réchauffer car $|Q_c| > |Q_f|$; la machine fournit plus à la source chaude qu'elle ne prend à la source froide, la différence provenant du travail.

► Remarque :

Un climatiseur fonctionne sur le même principe qu'une machine frigorifique.

d) Pompe à chaleur

Pour une pompe à chaleur, la source chaude est l'habitation à chauffer et la source froide peut être l'air extérieur ou une importante quantité d'eau (lac). Dans ce cas, la grandeur intéressante est le transfert thermique Q_c fourni à la source chaude et puisque $Q_c < 0$ et $W > 0$, l'efficacité est :

$$e = \left| \frac{Q_c}{W} \right| = -\frac{Q_c}{W} \quad (41.13)$$

soit avec $W = -(Q_c + Q_f)$:

$$e = \frac{Q_c}{Q_c + Q_f} = \frac{1}{1 + \frac{Q_f}{Q_c}} \quad (41.14)$$

Comme précédemment, on utilise l'inégalité de Clausius pour exprimer $\frac{Q_c}{Q_f}$:

$$\frac{Q_c}{T_c} + \frac{Q_f}{T_f} \leq 0 \Leftrightarrow \frac{Q_f}{T_f} \leq -\frac{Q_c}{T_c}$$

Or $Q_c < 0$ d'où :

$$\frac{Q_f}{Q_c} \geq -\frac{T_f}{T_c} \Leftrightarrow 1 + \frac{Q_f}{Q_c} \geq 1 - \frac{T_f}{T_c} (> 0)$$

D'où finalement :

$$e \leq \frac{1}{1 - \frac{T_f}{T_c}} = \frac{T_c}{T_c - T_f} \tag{41.15}$$

On trouve encore le même type de résultat, à savoir que l'efficacité d'une pompe à chaleur est inférieure à l'efficacité de Carnot. Il y a égalité dans le cas réversible. Les pompes à chaleur à condensation fonctionnent suivant le principe décrit par le schéma de la figure 41.5.

On effectue une application numérique, pour se rendre compte de ce que représente l'efficacité. La pièce à chauffer est à une température $T_c = 293\text{ K}$, tandis que la source froide est un lac à $T_f = 273\text{ K}$. Si on calcule l'efficacité, on trouve $e = 14,7$. Cela signifie que, pour un travail électrique donné $W_{\text{élec}}$, la source chaude recevra un transfert thermique $14,7\ W_{\text{élec}}$, alors que si on avait utilisé un convecteur avec une résistance chauffante, la source chaude n'aurait reçu que $W_{\text{élec}}$. On voit l'intérêt d'une pompe à chaleur. Alors pourquoi ne pas en utiliser systématiquement ? Tout simplement pour une raison de coût initial : une telle machine n'est amortie que pour une installation importante.

4.3 Tableau récapitulatif

On rappelle que l'indice (f) est relatif à la source froide et l'indice (c) à la source chaude.

	Moteur
Signe	$W < 0, \ Q_c > 0, \ Q_f < 0$
Grandeur intéressante	W
Grandeur coûteuse	Q_c
Rendement ou efficacité	$\rho = \frac{-W}{Q_c}$
Valeur de Carnot	$\rho_C = 1 - \frac{T_f}{T_c}$

	Pompe à chaleur	Machine frigorifique
Signe	$W > 0, \ Q_c < 0, \ Q_f > 0$	$W > 0, \ Q_c < 0, \ Q_f > 0$
Grandeur intéressante	Q_c	Q_f
Grandeur coûteuse	W	W
Rendement ou efficacité	$e = \frac{-Q_c}{W}$	$e = \frac{Q_f}{W}$
Valeur de Carnot	$e_C = \frac{1}{1 - \frac{T_f}{T_c}}$	$e_C = \frac{1}{\frac{T_c}{T_f} - 1}$

5. Autres exemples de machines thermiques

5.1 Machine avec écoulement de fluide

Dans de nombreuses machines, il faut tenir compte de l'écoulement du fluide, par exemple les turbines, les turbocompresseurs, les réacteurs... On généralise le raisonnement étudié lors de la détente de Joule-Thomson dans le chapitre sur les applications du premier principe.

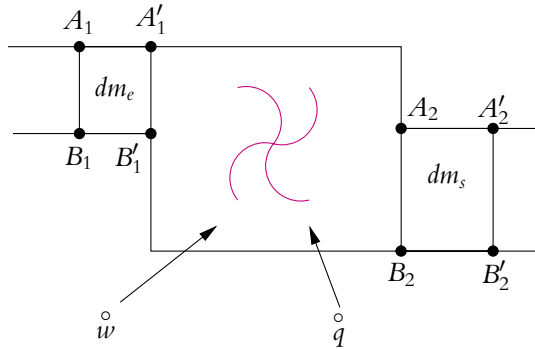


Figure 41.6 Schéma d'une machine avec écoulement.

- Le fluide s'écoule avec un débit massique D_m .
- Il reçoit de l'extérieur une puissance mécanique $\dot{w}$. Cette puissance est échangée avec la machine (turbine par exemple), ce n'est pas celle due aux forces de pression à gauche et à droite.
- Il reçoit aussi de l'extérieur une puissance thermique $\dot{q}$.
- Il subit une variation d'altitude $z_2 - z_1$.
- Il subit une variation de vitesse $c_2 - c_1$.

On sera amené à raisonner sur des grandeurs massiques et on appellera :

- u l'énergie interne massique,
- v le volume massique,
- h l'enthalpie massique.

Le premier système considéré (S') entre deux instants voisins t et $t + dt$ est la partie du fluide $A_1A_2B_1B_2$ qui devient $A'_1A'_2B'_1B'_2$. Étant donné que l'écoulement est stationnaire, les masses dm_e et dm_s des parties $A_1B_1A'_1B'_1$ et $A_2B_2A'_2B'_2$ sont égales et on note dm leur valeur commune.

Donc, la variation d'énergie totale entre les instants t et $t + dt$ est (voir expression (36.1)) :

$$dU_{(S')} + dE_{c(S')} + dEp_{(S')} = \delta W_{f,\text{pression}} + \overset{\circ}{w} dt + \overset{\circ}{q} dt \quad (41.16)$$

où $\delta W_{f,\text{pression}}$ est le travail des forces de pression. Comme dans l'étude de la détente de Joule-Thomson (Cf. expression (37.18)), il faut faire intervenir les volumes balayés. On note v_1 (respectivement v_2), les volumes massiques de $A_1B_1A'_1B'_1$ (respectivement $A_2B_2A'_2B'_2$). À gauche, le volume balayé est $v_1 dm_e$ et à droite $v_2 dm_s$, d'où le travail des forces de pression :

$$\delta W_{f,\text{pression}} = P_1 v_1 dm_e - P_2 v_2 dm_s = (P_1 v_1 - P_2 v_2) dm$$

On s'aperçoit que ce sont les grandeurs relatives aux volumes extrêmes qui interviennent. On est donc amené à définir le système fermé (S) de masse dm qui est en $A_1A'_1B_1B'_1$ initialement et en $A_2A'_2B_2B'_2$ finalement. Il a les caractéristiques suivantes :

État initial (1)	Altitude : z_1	État final (2)	Altitude : z_2
	Volume : $v_1 dm$		Volume : $v_2 dm$
	Pression : P_1		Pression : P_2
	Énergie interne :		Énergie interne :
	$dU_1 = u_1 dm$		$dU_2 = u_2 dm$
	Vitesse : c_1		Vitesse : c_2

Il s'agit maintenant de passer de (S') à (S). Sachant qu'en régime stationnaire, les énergies interne, cinétique et potentielle de la partie centrale $A'_1A_2B'_1B_2$ de (S') ne varient pas, on peut écrire :

$$\begin{aligned} dU_{(S')} &= dU_{(S)} \\ dE_{c(S')} &= dE_{c(S)} \\ dEp_{(S')} &= dEp_{(S)} \end{aligned} \quad (41.17)$$

► **Remarque :** Même si ce ne sont pas les mêmes molécules de gaz qui occupent les positions $A_1A'_1B_1B'_1$ et $A_2A'_2B_2B'_2$ dans les systèmes (S) et (S') aux mêmes instants, les tranches de gaz ont les mêmes caractéristiques puisqu'elles sont aux mêmes positions et que l'écoulement est stationnaire.

En combinant les expressions (41.16) et (41.17), on obtient :

$$dU_{(S)} + dE_{c(S)} + dEp_{(S)} = \delta W_{f,\text{pression}} + \overset{\circ}{w} dt + \overset{\circ}{q} dt$$

soit :

$$\left(u_2 + gz_2 + \frac{1}{2}c_2^2\right) dm - \left(u_1 + gz_1 + \frac{1}{2}c_1^2\right) dm = (P_1 v_1 - P_2 v_2) dm + \dot{w} dt + \dot{q} dt$$

Sachant que $D_m = \frac{dm}{dt}$, on obtient :

$$\left(u_2 + P_2 v_2 + gz_2 + \frac{1}{2}c_2^2\right) D_m - \left(u_1 + P_1 v_1 + gz_1 + \frac{1}{2}c_1^2\right) D_m = \dot{w} + \dot{q}$$

soit :

$$\begin{aligned} D_m \left(h_2 + gz_2 + \frac{1}{2}c_2^2\right) - D_m \left(h_1 + gz_1 + \frac{1}{2}c_1^2\right) &= \dot{w} + \dot{q} \\ D_m \Delta \left(h + gz + \frac{1}{2}c^2\right) &= \dot{w} + \dot{q} \end{aligned} \quad (41.18)$$

où $h = u + Pv$ est l'enthalpie massique.

Exemple

On pompe l'eau d'un puits à 30 m de profondeur, à la température 363 K, avec un débit volumique de $D_v = 120 \text{ L.min}^{-1}$. Cette eau est récupérée dans un réservoir à la température 273 K après être passée dans un échangeur thermique où elle reçoit une puissance thermique $\dot{q}$. La puissance mécanique de la pompe est $\dot{w} = 2 \text{ kW}$. On cherche à calculer $\dot{q}$ sachant que la capacité thermique de l'eau est $C = 4,2 \text{ kJ.kg}^{-1}$ et que l'on néglige son énergie cinétique. L'accélération de la pesanteur vaut $9,81 \text{ ms}^{-2}$.

Il suffit d'appliquer l'équation (41.18) avec ici

$$c_1 = c_2 = 0, \quad h_2 - h_1 = C(T_2 - T_1) = -376,2 \text{ kJ}$$

et $D_m = D_v \rho_{\text{eau}} = \frac{120 \cdot 10^{-3}}{60} \times 10^3 = 2 \text{ kg.s}^{-1}$ où ρ_{eau} est la masse volumique de l'eau. On trouve, en arrondissant, $\dot{q} = -754 \text{ kJ.s}^{-1}$. L'eau pompée cède du transfert thermique et peut servir de source d'énergie (géothermie).

5.2 Cas où le cycle n'est pas fermé

Dans certaines machines comme des turbocompresseurs (avions, marteau-piqueur), le même fluide ne subit pas un cycle entier. Par exemple dans un réacteur d'avion :

- l'air est admis à (P_1, T_1) ,
- AB : l'air est comprimé adiabatiquement à (P_2, T_2) ,
- BC : l'air est brûlé avec le carburant jusqu'à l'état final (P_2, T_3) ,
- CD les gaz brûlés subissent une détente dans la turbine jusqu'à (P_4, T_4) ,
- DE : à la sortie de la tuyère, les gaz sont éjectés dans l'atmosphère sans être revenus à (P_1, T_1) .

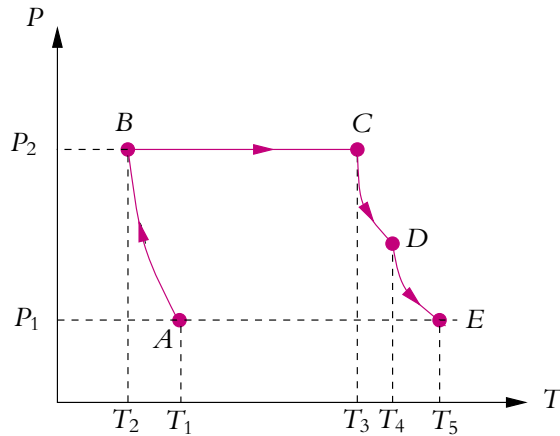


Figure 41.7 Turboréacteur.

Puisqu'il ne s'agit pas d'un cycle fermé (Cf. figure 41.7), on ne peut plus écrire $\Delta U = 0$ ou $\Delta S = 0$.

5.3 Moteur à essence

Les moteurs à essence fonctionnent suivant un cycle théorique proposé par le physicien français Beau de Rochas en 1862. Le moteur fut réalisé par l'allemand Otto une quinzaine d'années plus tard. Ces moteurs, qu'ils soient à deux temps (scooter, outils de jardinage) ou à quatre temps (moteur de voiture), sont dits à *moteurs à explosion* car il est nécessaire de produire une étincelle à l'aide d'une bougie pour provoquer l'inflammation du mélange air-carburant (contrairement au moteur diesel).

Dans un premier temps, on décrit le moteur à quatre temps. On a représenté sur la figure 41.8 le cycle théorique et sur la figure 41.9 le cycle réel qui, comme on le voit, se rapproche du cycle théorique.

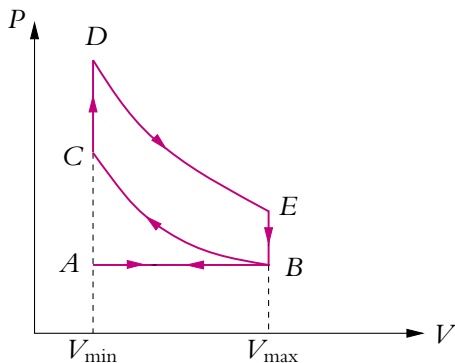


Figure 41.8 Cycle théorique du moteur à quatre temps.

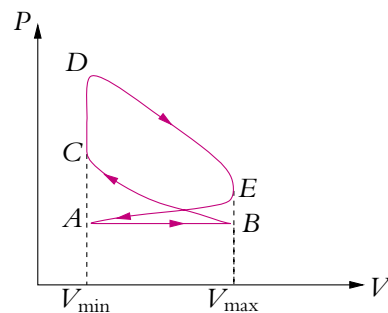


Figure 41.9 Cycle réel du moteur à quatre temps.

En A le piston est en bout de course et le cylindre offre le volume minimal $V_{\min}$. L'évolution est la suivante :

- 1^{er} temps (AB) : admission, la soupape d'admission est ouverte, le piston descend en aspirant le mélange air-carburant jusqu'au volume $V_{\max}$.
- 2^e temps (BCD) : le piston remonte et comprime le gaz adiabatiquement jusqu'en C puis l'étincelle est produite par la bougie provoquant la combustion. La pression augmente très rapidement mais le piston n'a pas le temps de bouger (évolution isochore CD).
- 3^e temps (DE) : les gaz brûlés sous forte pression repoussent le piston. C'est une détente adiabatique avec production de travail.
- 4^e temps (EBA) : la soupape d'échappement s'ouvre, provoquant une rapide baisse de pression isochore (EB), puis le piston remonte pour refouler les gaz brûlés (BA).

Les mouvements du piston sont représentés sur les figures 41.10 à 41.13.

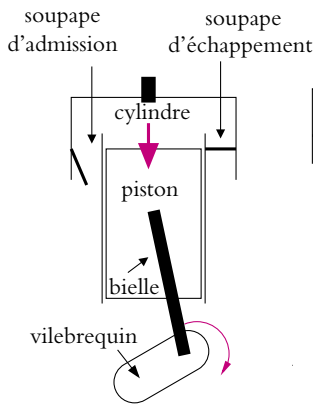


Figure 41.10 Admission.

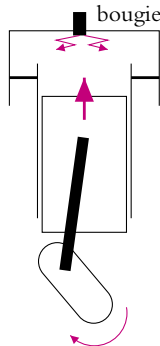


Figure 41.11 Compression puis explosion.

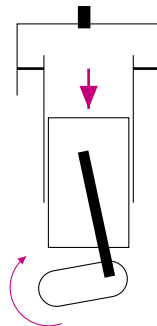


Figure 41.12 Détente avec production de travail.

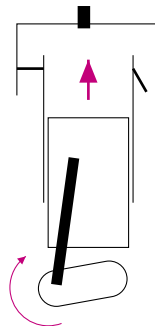


Figure 41.13 Échappement.

On s'aperçoit donc que, dans un moteur à quatre temps, le piston fait deux allers et retours pour décrire un cycle. Dans un moteur de voiture, les différents cylindres fonctionnent avec un décalage de manière à ce que le piston remonte dans certains cylindres quand il descend dans d'autres. Cela est dû au fait que les bielles des différents cylindres ne sont pas fixées toutes du même côté du vilebrequin.

Le lecteur pourra s'exercer à calculer le rendement de ce moteur. Il s'exprime uniquement en fonction du taux de compression, c'est-à-dire $\frac{V_{\max}}{V_{\min}}$. Si on considère le gaz parfait et les adiabatiques réversibles, on trouve :

$$\rho = 1 - \left(\frac{V_{\max}}{V_{\min}} \right)^{(1-\gamma)} \quad (41.19)$$

Comme $1 - \gamma < 0$, l'expression précédente montre que ρ est d'autant plus grand que le taux de compression est important. Ce paramètre est pris en compte dans la fabrication des carburants qui doivent supporter des forts taux de compression et ne pas exploser avant l'étincelle de la bougie.

Qu'est-ce qui différencie un moteur deux temps d'un moteur quatre temps ? Les quatre phases, réparties sur deux tours dans un moteur quatre temps, ne sont plus réparties que sur un tour de vilebrequin dans le moteur à deux temps. Le diagramme est semblable à celui de la figure 41.8 sans la partie horizontale (AB). L'admission et le refoulement se font en (DE) spontanément. En fait, la légère surpression en E permet l'expulsion spontanée des gaz brûlés quand la soupape d'échappement s'ouvre. Les quatre temps se réduisent alors à deux temps :

- 1^{er} temps (BCD) : compression adiabatique du mélange air-carburant jusqu'en C puis combustion (évolution isochore CD).
- 2^e temps (DEB) : détente adiabatique avec production de travail (DE) puis échappement et admission d'un nouveau mélange air - carburant (EB).

Le vilebrequin du moteur à deux temps n'effectuant qu'un tour au lieu de deux, ce moteur devrait être deux fois plus puissant que le moteur à quatre temps. Cependant, il fonctionne dans des conditions moins favorables et on trouve un rapport environ 1,4 entre les puissances.

5.4 Moteur Diesel

Comme on l'a signalé dans le paragraphe précédent, le moteur Diesel ne nécessite pas de bougie. Il suit le cycle théorique représenté sur la figure 41.14, le cycle réel étant représenté sur la figure 41.15.

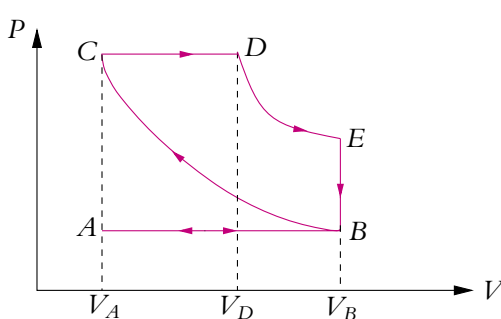


Figure 41.14 Cycle théorique du moteur Diesel.

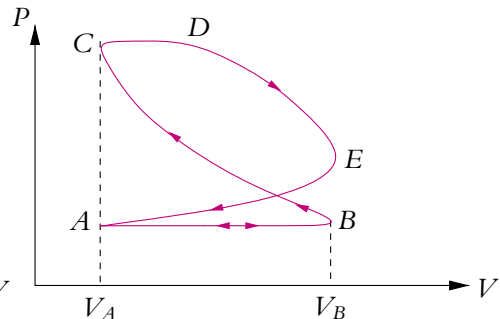


Figure 41.15 Cycle réel du moteur Diesel.

Le fonctionnement se décompose de la manière suivante :

- 1^{er} temps (AB) : admission de l'air ;
- 2^{ème} temps (BC) : compression adiabatique de l'air jusqu'à une pression élevée : l'air atteint une température élevée de l'ordre de $600\text{ }^{\circ}\text{C}$;

- 3^{ème} temps (CD) : injection du carburant qui s'enflamme spontanément au contact de l'air chaud, le taux d'injection du carburant est réglé de manière à ce que l'évolution soit isobare ;
- 4^{ème} temps (DE) : lorsque la combustion est terminée, on laisse le mélange se détendre adiabatiquement en repoussant le piston, il y a production de travail ;
- 5^{ème} temps (EB) : la soupape d'échappement s'ouvre provoquant une rapide baisse de pression ;
- 6^{ème} temps (BA) : le piston refoule les gaz brûlés.

Si on définit les rapports $a = \frac{V_B}{V_C}$ et $b = \frac{V_E}{V_D}$, en supposant les transformations réversibles et les gaz parfaits, le rendement vaut :

$$\rho = 1 - \frac{1}{\gamma} \left(\frac{a^{-\gamma} - b^{-\gamma}}{a^{-1} - b^{-1}} \right). \quad (41.20)$$

5.5 Cycle de Rankine

De nombreuses turbines fonctionnent grâce à un cycle diphasé liquide-vapeur, en particulier les centrales de production d'électricité. Les centrales nucléaires à eau pressurisée fonctionnent suivant le cycle de Rankine qui est décrit dans la suite.

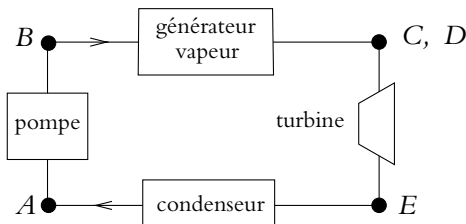


Figure 41.16 Schéma de l'installation.

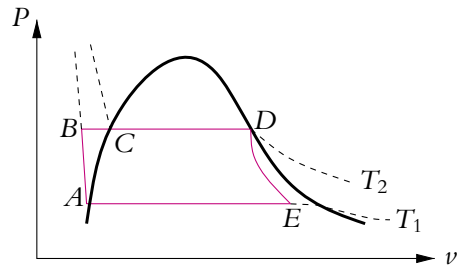


Figure 41.17 Diagramme de Clapeyron.

Le système est constitué de l'eau qui circule dans le circuit secondaire dans le cas d'une centrale nucléaire, le circuit primaire étant celui en contact avec le cœur du réacteur. On conseille au lecteur de suivre la série des transformations simultanément sur les deux schémas des figures 41.16 et 41.17.

- En A, l'eau qui sort du condenseur est liquide.
- De A à B, l'eau subit une compression adiabatique supposée réversible dans la pompe. Si on suppose l'eau incompressible, $cdT = TdS$, si bien que sa température ne varie quasiment pas et que les points A et B sont sur une isotherme.
- L'eau passe dans le générateur de vapeur : sa pression est 70 bars et sa température est 290 °C. C'est un échangeur qui permet les échanges thermiques entre le circuit

primaire, dont l'eau est chauffée par les réactions nucléaires du coeur, et le circuit secondaire. On peut décomposer en deux transformations : de B à C échauffement de l'eau liquide de manière isobare puis de C à D vaporisation complète de l'eau liquide.

- De D à E , la vapeur d'eau passe dans une turbine où elle subit une détente adiabatique supposée réversible en fournissant du travail. Durant cette détente, une fraction de la vapeur redevient liquide.
- De E à A , le mélange liquide-vapeur passe dans un condenseur (tour de refroidissement), dans lequel il achève de se liquéfier. À la sortie, la pression de l'eau est 50 mbar et la température 32 °C.

Un des défauts de ce type de machine réside dans l'apparition de gouttelettes d'eau liquide lors de la détente de la turbine, ce qui provoque une corrosion accélérée des ailettes. Pour éviter cela, il faut modifier le cycle pour que le point E ne se trouve pas dans la zone liquide vapeur. C'est le choix fait par Hirn et décrit dans le paragraphe suivant.

5.6 Cycle de Hirn

Entre le générateur de vapeur et la turbine, on ajoute un surchauffeur, dont le rôle est de chauffer la vapeur de manière isobare jusqu'à la température maximale compatible avec la résistance des ailettes de la turbine (Cf. figure 41.18). Dans ce cas, le cycle est modifié comme cela apparaît sur la figure 41.19.

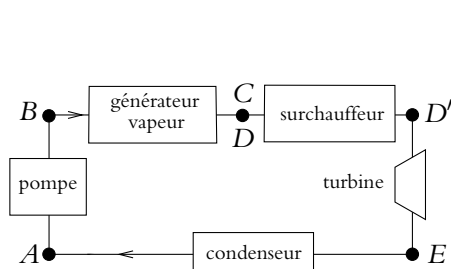


Figure 41.18 Schéma de l'installation.

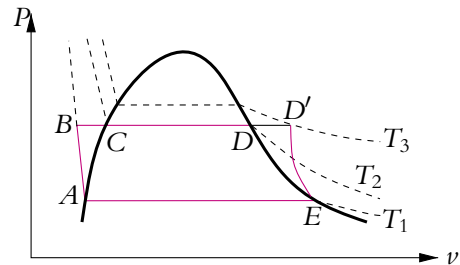


Figure 41.19 Diagramme de Clapeyron.

Un deuxième avantage du cycle de Hirn est d'accroître la puissance de la machine. Il arrive même qu'on effectue une deuxième surchauffe pour augmenter encore le rendement.

A. Applications directes du cours

1. Moteur réel

Un moteur réel fonctionnant entre deux sources de chaleur, l'une à $T_F = 400$ K, l'autre à $T_C = 650$ K, produit 500 J par cycle, pour 1500 J de transfert thermique fourni.

1. Comparer son rendement à celui d'une machine de Carnot fonctionnant entre les deux mêmes sources.
2. Calculer l'entropie créée par cycle, notée S_{cr} .
3. Montrer que la différence entre le travail fourni par la machine de Carnot et la machine réelle est égale à $T_F S_{cr}$, pour une dépense identique.

2. Pompe à chaleur

Une pompe à chaleur quasi-statique fonctionne entre deux sources constituées l'une par l'eau d'un lac à température constante $t = 5^\circ\text{C}$, l'autre par une masse M d'eau thermiquement isolée ($c = 4,18 \text{ J}\cdot\text{g}^{-1}\cdot\text{K}^{-1}$). La machine fonctionne dans un sens tel que la masse M d'eau s'échauffe.

1. Quelle est la relation entre la température t de l'eau et l'énergie W fournie à la machine ?
2. Calculer W pour une tonne d'eau quand elle a atteint $t_f = 40^\circ\text{C}$, partant d'une température initiale $t_i = 10^\circ\text{C}$.
3. Quelle aurait été l'élévation de température si la même énergie avait été utilisée directement dans une résistance chauffante ?

3. Moteur thermique, d'après ENAC 2004

Un moteur thermique fonctionne de manière réversible entre deux sources dont les températures T_c et T_f peuvent évoluer au cours du temps à cause des échanges thermiques avec la machine. La source froide est constituée d'une masse $M = 100$ kg d'eau en totalité à l'état de glace fondante à la température $T_{f,0} = 273$ K. La source chaude est constituée d'une masse $2M$ d'eau liquide à la température $T_{c,0} = 373$ K. On donne la capacité thermique massique de l'eau liquide $C = 4,18 \text{ kJ}\cdot\text{kg}^{-1}\cdot\text{K}^{-1}$ et la chaleur latente massique de fusion de la glace $L = 335,6 \text{ kJ}\cdot\text{kg}^{-1}$ à la température $T_{f,0}$.

1. Déterminer la température $T_{c,1}$ de la source chaude quand la totalité de la glace de la source froide a fondu.
 1. $T_{c,1} = 322$ K,
 2. $T_{c,1} = 300$ K,
 3. $T_{c,1} = 305$ K,
 4. $T_{c,1} = 352$ K.
2. Calculer le travail W_1 fourni alors par le moteur.
 1. $W_1 = -7,03$ MJ,
 2. $W_1 = -9,08$ MJ,
 3. $W_1 = -3,86$ MJ,

4. $W_1 = -3,14 \text{ MJ}$.
3. Le moteur s'arrête lorsque les deux sources sont à la même température T_0 .
 1. $T_0 = 290 \text{ K}$,
 2. $T_0 = 325 \text{ K}$,
 3. $T_0 = 313 \text{ K}$,
 4. $T_0 = 305 \text{ K}$.
4. Calculer le rendement thermique global du moteur.
 1. $\eta = 0,32$,
 2. $\eta = 0,25$,
 3. $\eta = 0,72$,
 4. $\eta = 0,18$.
5. Même question si les températures des deux sources avaient été maintenues constantes.
 1. $\eta_0 = 0,27$,
 2. $\eta_0 = 0,32$,
 3. $\eta_0 = 0,49$,
 4. $\eta_0 = 0,78$.

B. Exercices et problèmes

1. Pompe à chaleur, d'après ENSTIM 2005

Une pompe à chaleur effectue le cycle de Joule inversé suivant. L'air pris dans l'état A de température T_0 et de pression P_0 est comprimé suivant une adiabatique quasistatique jusqu'au point B où il atteint la pression P_1 . Le gaz est ensuite refroidi à pression constante et atteint la température finale de la source chaude T_1 correspondant à l'état C . Le gaz est encore refroidi dans une turbine suivant une détente adiabatique quasistatique pour atteindre l'état D de pression P_0 . Le gaz se réchauffe enfin à pression constante au contact de la source froide et retrouve son état initial.

On considère que l'air est un gaz parfait de coefficient isentropique $\gamma = 1,4$. On pose $\beta = 1 - \frac{1}{\gamma}$ et $a = \frac{P_1}{P_0}$. On prendra $T_0 = 283 \text{ K}$, $T_1 = 298 \text{ K}$, $a = 5$ et $R = 8,31 \text{ J.K}^{-1}.\text{mol}^{-1}$.

1. Représenter le cycle parcouru par le gaz dans le diagramme de Clapeyron donnant la pression en fonction du volume.
2. Rappeler les conditions nécessaires pour assurer la validité des formules de Laplace. Donner la formule de Laplace relative à la pression et à la température.
3. En déduire l'expression des températures T_B et T_D des états B et D en fonction de T_0 , T_1 , a et β . Préciser leur valeur numérique.
4. Définir l'efficacité e de la pompe à chaleur en fonction des transferts thermiques du cycle.
5. En déduire l'expression de e en fonction de a et β . Donner sa valeur numérique.

6. Quelles doivent être les transformations du gaz si on fait fonctionner la pompe à chaleur suivant un cycle de Carnot réversible entre les températures T_0 et T_1 ?
7. Établir l'expression de son efficacité e_r . Donner sa valeur numérique.
8. Comparer e et e_r . Proposer une explication à ce résultat.
9. Déterminer l'expression de l'entropie créée S_c pour une mole d'air au cours du cycle de Joule en fonction de R , β et $x = a^\beta \frac{T_0}{T_1}$.
10. Étudier le signe de S_c en fonction de x . Était-ce prévisible ?
11. Calculer sa valeur ici.
12. Sachant qu'en régime permanent, les fuites thermiques s'élèvent à $Q_f = 20$ kW, calculer la puissance du couple compresseur – turbine qui permet de maintenir la température de la maison constante.

2. Mise à l'équilibre thermique

Deux sources liquides de même capacité thermique C , initialement à la température T_1 et T_2 , sont mises en contact, mais isolées de l'extérieur. On appelle σ la variation d'entropie.

1. Calculer σ .
2. Calculer le travail nécessaire pour ramener les sources à leur température initiale. On dispose pour cela d'une machine thermique fonctionnant entre les deux sources précédentes et un thermostat idéal à la température T_0 .
3. Montrer que le thermostat à T_0 est nécessaire.

3. Optimisation du chauffage d'un local, d'après Centrale 1984

On souhaite maintenir la température d'un local à $\theta_1 = 20^\circ\text{C}$ alors que la température extérieure est $\theta_2 = -2^\circ\text{C}$. L'énergie thermique nécessaire est de 32 MJ par heure. Par un système de chauffage central, on brûle a litres de fuel par jours (a est de l'ordre de 25 litres, sa valeur exacte n'intervient pas dans l'exercice).

1. On utilise une pompe à chaleur fonctionnant réversiblement entre le local et l'extérieur, le fuel servant à faire fonctionner le moteur de cette pompe, comme nous le décrit la question 2. Calculer l'efficacité de la pompe à chaleur. En déduire la puissance consommée par le moteur de cette pompe.
2. Deux conseillers proposent chacun un dispositif qu'ils déclarent thermodynamiquement plus avantageux que le chauffage central :
 - dispositif du conseiller n°1 : les a litres de fuel sont brûlés, l'énergie thermique Q récupérée permet d'assurer la vaporisation de l'eau d'une chaudière auxiliaire à la température $\theta_3 = 210^\circ\text{C}$ qui sert de source chaude à un moteur ditherme réversible dont la source froide est le local, le travail fourni servant à faire fonctionner la pompe à chaleur étudiée à la question 1 ;
 - dispositif du conseiller n°2 : le principe est le même mais la chaudière auxiliaire est à la température $\theta_4 = 260^\circ\text{C}$ et le moteur fonctionne entre cette chaudière et l'air extérieur.

Dans les deux cas, on suppose que toute l'énergie thermique Q obtenue par combustion du fuel est fournie par la chaudière auxiliaire au fluide du moteur.

- Dans les deux cas, calculer la durée Δt (en jours) pendant lequel le chauffage sera assuré avec les a litres de fuel.
- Lequel de ces deux systèmes est le plus économique ? Le système le plus économique tire-t-il son avantage de la température à laquelle fonctionne la chaudière auxiliaire ou cet avantage se maintient-il pour $\theta_3 = \theta_4$? Expliquer.
- La durée de chauffage, pour une quantité de fuel donnée, augmentant avec la température de la chaudière auxiliaire, un troisième conseiller prétend qu'il pourra, en utilisant le même dispositif, augmenter la durée de chauffage. A-t-il tort ou raison ? Pourquoi ?

4. Étude de la turbine d'un réacteur à eau pressurisée (REP), d'après ENSTIM 1996

Le parc de production nucléaire français est composé de centrales de la filière REP :

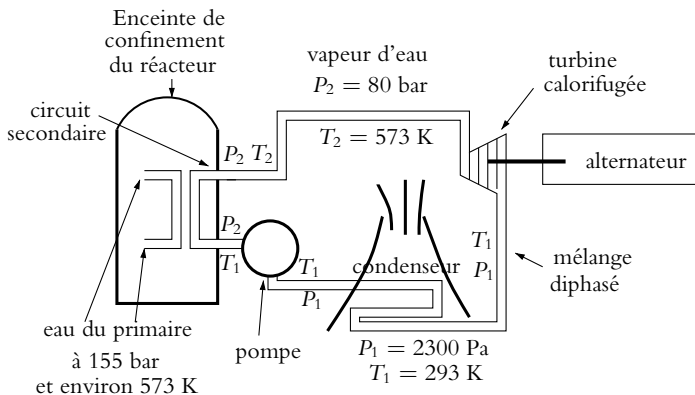


Figure 41.20

On étudie l'eau ($M = 18 \text{ g.mol}^{-1}$) en circuit fermé du circuit secondaire. On propose de modéliser son évolution au prix de quelques approximations par le cycle suivant :

- État A : l'eau qui sort du condenseur est liquide sous la pression P_1 et à la température T_1 .
- Évolution AB : elle subit dans la pompe une compression durant laquelle sa température ne varie pratiquement pas. On considérera que les échanges thermiques sont négligeables lors de cette compression qui l'amène dans l'état B sous la pression P_2 et à la température T_1 .
- Évolution BD : elle passe ensuite dans un échangeur qui permet les transferts thermiques entre le circuit primaire et le circuit secondaire. On peut décomposer en deux transformations ce qui se passe alors :
 - l'eau liquide s'échauffe de manière isobare (sous la pression P_2), jusqu'à l'état C (pression P_2 , température T_2) ;
 - l'eau liquide se vaporise entièrement, jusqu'à l'état D sous la pression P_2 et à la température T_2 .

- Évolution DE : la vapeur d'eau se détend de manière réversible dans une turbine calorifugée jusqu'à la pression P_1 et la température T_1 (état E). Durant cette détente, une fraction $(1-x)$ de l'eau redevient liquide, et x reste gazeuse : x est donc le titre en vapeur.
- Évolution EA : la vapeur restant se condense à la température T_1 .

L'écoulement est stationnaire.

Dans le tableau suivant, on donne : la pression de vapeur saturante P_s en bar, les volumes massiques v_l du liquide et v_v de la vapeur en $\text{m}^3.\text{kg}^{-1}$, les enthalpies massiques h_l du liquide et h_v de la vapeur en $\text{kJ}.\text{kg}^{-1}$ et enfin les entropies massiques s_l du liquide et s_v de la vapeur en $\text{kJ}.\text{K}^{-1}.\text{kg}^{-1}$.

$T(\text{K})$	P_s	v_l	v_v	h_l	h_v	s_l	s_v
293	0,023 (P_1)			85	2540	0,3	8,7
573	80 (P_2)	1, 31.10 ⁻³	0,026	1290	2890		6,0

On rappelle la valeur de la constante des gaz parfaits : $R = 8,314 \text{ J}.\text{K}^{-1}.\text{mol}^{-1}$.

1. a) Quelle est la variation d'entropie d'un corps pur et sa variation d'enthalpie lors d'une vaporisation totale ?
b) Calculer l'entropie massique de l'eau liquide, $s_l(573)$, à la température de 573 K et sous une pression de 80 bar.
c) Sachant que la vapeur d'eau sous une pression de 0,023 bar et à la température de 293 K peut être considérée comme un gaz parfait, calculer son volume massique $v_v(293)$.
2. Tracer le cycle de l'eau sur un diagramme de Clapeyron (P, v) en plaçant les points correspondant aux états A, B, C, D et E.
3. a) Démontrer que la transformation DE est isentropique.
b) Calculer le titre en vapeur dans l'état E.
4. On notera W_a le travail reçu par l'alternateur par unité de masse du fluide écoulé (on négligera tout frottement). La turbine est calorifugée et horizontale. Calculer la valeur numérique de W_a .
5. a) Pour l'eau liquide, la capacité thermique massique est $c = 4,2 \text{ kJ}.\text{K}^{-1}.\text{kg}^{-1}$. Calculer le transfert thermique (par unité de masse de fluide écoulé) du système avec l'échangeur du circuit primaire Q_{BD} .
b) On définit l'efficacité $e = \frac{W_a}{Q_{BD}}$. La calculer. Que néglige-t-on dans cette définition ?
c) Calculer l'efficacité maximale qu'on aurait pu avoir avec les mêmes sources. Quelle conclusion peut-on en tirer ?

5. Procédé Charles de liquéfaction de l'azote

On donne (les indices correspondent aux numéros des états décrits ci-dessous) :

- Températures : $T_1 = 122 \text{ K}$, $T_2 = 290 \text{ K}$, $T_3 = 160 \text{ K}$, $T_4 = 78 \text{ K}$.
- Entropies massiques (en $\text{kJ}.\text{K}^{-1}.\text{kg}^{-1}$) : $s_1 = 4,40$, $s_2 = 2,68$, $s_3 = 1,76$, $s_A = 0,44$, $s_B = 2,96$.
- Enthalpies massiques (en $\text{kJ}.\text{kg}^{-1}$) : $h_2 = 421$, $h_3 = 214$, $h_6 = 34$.

L'allure du diagramme de Clapeyron (P, ν), où ν est le volume massique, est donnée sur la figure ci-contre, sur laquelle on a fait figurer les isothermes à T_2, T_3 et T_4 .

Le diazote subit à partir du point 1 les transformations suivantes :

- une compression isotherme de 1 à 2 dans le compresseur supposée réversible ;
- la traversée d'un échangeur thermique de 2 à 3 avec évolution isobare ($P_3 = 200$ bar) jusqu'à la température T_3 ;
- une détente adiabatique réversible de 3 à 4 dans une turbine ; on obtient en 4 un mélange liquide-vapeur à la température T_4 ;
- le diazote passe alors dans un séparateur (évolution isotherme), pour récupérer en 5 la vapeur seule et en 6 le liquide seul ;
- 5-1 : la vapeur passe dans l'échangeur thermique, en croisant le diazote évoluant de 2 à 3, l'évolution est isobare.

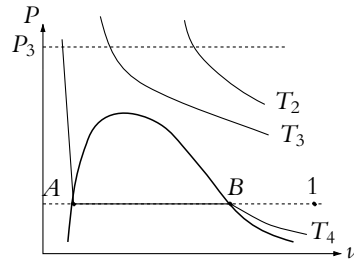


Figure 41.21

1. Représenter la suite des transformations sur un diagramme de Clapeyron en tenant compte des données de l'énoncé.
2. Calculer l'enthalpie massique de vaporisation L_v à 78 K.
3. Calculer le titre massique de vapeur x_4 obtenue en 4.
4. Exprimer l'enthalpie massique en 4 en fonction de h_5, h_6 et x_4 .
- 5.

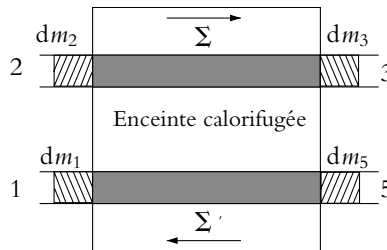


Figure 41.22

Un échangeur thermique est un appareil dans lequel se croisent deux tuyaux transportant les fluides, l'ensemble étant calorifugé. On considère le système (S) constitué à l'instant t des systèmes Σ et Σ' ainsi que des masses dm_2 et dm_5 qui entrent dans l'échangeur entre t et $t + dt$. À l'instant $t + dt$, ce système est constitué de Σ , de Σ' et des masses dm_1 et dm_3 qui sont sorties de l'échangeur entre t et $t + dt$. Montrer qu'en régime stationnaire :

$$h_3 - h_2 = x_4 (h_5 - h_1)$$

En déduire h_1 et sa valeur numérique.

6. Calculer le transfert thermique reçu par 1 kilogramme de diazote lors de la compression 1-2.

7. On rappelle que dans le cas d'un fluide en écoulement, le premier principe prend la forme suivante :

$$\Delta H = W_{\text{autre}} + Q$$

où W_{autre} est le travail autre que celui des forces de pression.

Calculer le travail reçu de la part du compresseur par 1 kilogramme de diazote lors de la compression 1-2. En déduire la masse de diazote liquide produite en une heure pour une machine de puissance 100 kW.

8. Quelle est la puissance récupérée par la turbine ?

6. Moteur Diesel, d'après ESIM MP, PSI 2003

Dans l'étude théorique du moteur Diesel, on suppose que le fonctionnement est décrit par le cycle de Diesel dans lequel la combustion est supposée s'effectuer à pression constante. Dans les moteurs actuels et notamment dans ceux dits rapides, on a cherché à réaliser une combustion procédant sous deux régimes : la combustion commence à volume constant et se termine à pression constante. Ce cycle de combustion est baptisé différemment suivant les pays (Seiliger en Allemagne, Trinkler en Russie, Sabathé en France). On peut également le qualifier de cycle mixte. Il consiste en une compression adiabatique AB , une détente isochore BC , un échauffement isobare CD , une détente adiabatique DE et une compression isochore EA .

L'objet de ce problème est d'étudier un tel moteur. On considère un moteur Diesel à six cylindres fonctionnant suivant le cycle mixte à quatre temps. Le moteur à quatre temps effectue un cycle tous les deux tours. Il présente un rapport de compression volumétrique $a = \frac{V_A}{V_B} = 15$ et une cylindrée $B = V_A - V_B = 15$ L. V_A est le volume maximal d'air aspiré, V_B le volume disponible dans le cylindre au moment où commence l'injection du combustible, V_D le volume quand la combustion se termine. On note $\epsilon = \frac{V_D}{V_B}$ le rapport d'injection.

On désigne par P_A la pression de l'air à l'aspiration, P_B celle au moment où commence l'injection du combustible et P_C la pression maximale atteinte dans le cylindre. On suppose que le cycle est décrit de manière quasistatique.

On admet que le gaz constitué essentiellement d'air est assimilable à un gaz parfait et que la masse de carburant est négligeable devant celle de l'air pour un cycle. On raisonnera sur un cycle fictif fermé sans se préoccuper des étapes consistant à évacuer les gaz issus de la combustion (échappement) pour les remplacer par de l'air frais (admission) qu'on modéliser par l'évolution isochore EA .

On note C_V et $C_P = 1,02 \text{ kJ.kg}^{-1}.\text{K}^{-1}$ les capacités thermiques massiques respectivement à volume et à pression constants, on admet qu'ils sont constants. Enfin $\gamma = \frac{C_P}{C_V} = 1,39$ et m désigne la masse d'air aspirée par cycle et par cylindre.

1. Sachant qu'un cycle est moteur lorsqu'il fournit du travail, dans quel sens le cycle est-il décrit ?

2. Exprimer les transferts thermiques lors des transformations BC , CD et EA en fonction des températures, de m et des capacités thermiques massiques.

3. Justifier qu'on définit le rendement par $r_{th} = -\frac{W}{Q_{BC} + Q_{CD}}$ sachant qu'un rendement est le rapport de ce qui est utile sur ce qu'on dépense.

4. Donner son expression en fonction des températures et de γ .
5. Exprimer les températures T_B , T_C , T_D et T_E en fonction de T_A , a , c , $d = \frac{P_D}{P_B}$ et γ .
6. En déduire l'expression du rendement en fonction de a , c , d et γ .
7. La combustion d'une masse m' de carburant produit un transfert thermique $\eta m'$ où $\eta = 42,2 \text{ MJ.kg}^{-1}$ désigne le pouvoir calorifique dont l'unité est J.kg^{-1} . On admet qu'il s'agit d'une constante indépendante des conditions durant la combustion. La masse de combustible pulvérisé dans un cylindre à chaque cycle est égale à αm avec $\alpha = 40.10^3 \text{ kg}$ par kilogramme d'air. Au début de la compression, la température est $T_A = 338 \text{ K}$ et la pression P_A . On désigne par μ la masse totale de combustible injectée par cycle et par cylindre et par μ_V celle qui brûle à volume constant. Le cycle est caractérisé par $k_V = 100 \frac{\mu_V}{\mu}$. On note $M_{\text{air}} = 29 \text{ g.mol}^{-1}$ la masse molaire de l'air.
Exprimer la masse d'air contenue dans le cylindre m en fonction de a , B , P_A , T_A , R et M_{air} .
8. Exprimer Q_{BC} en fonction de k_V , η , m et α .
9. En comparant les deux expressions obtenues pour Q_{BC} , établir l'expression de T_C et P_C en fonction de T_A , P_A , a , γ , k_V , η , α et C_V .
10. Exprimer Q_{CD} en fonction de k_V , η , m et α .
11. En déduire l'expression de T_D en fonction de T_A , a , γ , k_V , η , α et C_V .
12. Quelle valeur doit-on donner à k_V pour que la pression maximale dans le cylindre ne dépasse pas 65 bars sachant que $P_A = 1,0 \text{ bar}$?
13. Si k_V prend cette valeur, déterminer les valeurs de V_A , V_B , T_B , T_C , T_D , c , d , T_E et b .
14. Déterminer la valeur du rendement théorique.
15. Calculer la puissance du moteur sachant qu'il tourne à 2300 tours par minute.
16. La puissance réelle est-elle inférieure ou supérieure à cette valeur théorique? Justifier votre réponse.

The background is a solid light pink color. In the upper right quadrant, there is a bright, glowing white light source. From this source, several concentric, semi-transparent pink circles radiate outwards, creating a ripple effect across the page. The circles vary in opacity, with the ones closer to the light source being more vibrant and those further away being more faded.

Électromagnétisme

Ce chapitre est consacré à la description de la répartition des charges. On introduit la notion de distribution de charges puis on étudie les propriétés d'invariance de celles-ci ainsi que leurs conséquences.

1. Distributions de charges

La notion de charge électrique a été définie au chapitre 1. On peut rappeler ici quelques propriétés utiles pour l'étude de l'électromagnétisme. Les phénomènes d'attraction et de répulsion électrique peuvent s'expliquer par la notion de charge électrique. Celle-ci peut être de deux types : positive ou négative, elle est extensive, conservative, invariante par changement de référentiel et quantifiée.

Cette notion de charge électrique ne s'applique qu'à des charges ponctuelles. Or à l'échelle macroscopique, même un très petit échantillon contient un très grand nombre de charges élémentaires. Par exemple, un millimètre cube de dioxygène dans les conditions normales de température et de pression contient $8,6 \cdot 10^{17}$ charges élémentaires, la moitié est négative, l'autre positive pour que l'électroneutralité de la matière soit assurée.

L'électrisation de la matière déplace les charges : on a donc un surplus de charges d'un type en un endroit et d'un autre ailleurs. Ce surplus ΔN est très faible devant le nombre total de charges mais très grand devant 1 :

$$1 \ll \Delta N \ll N$$

On va donc, comme en thermodynamique, s'intéresser à des grandeurs moyennes, ce qui conduit naturellement à considérer des charges « continues » bien que la charge

électrique soit par nature quantifiée. Tout sera donc question d'échelle. On va préciser sur cet exemple les notions d'échelles microscopique, mésoscopique et macroscopique déjà abordées dans le cours de thermodynamique.

1.1 Échelles microscopique, mésoscopique et macroscopique

La notion de charge électrique introduite au paragraphe précédent s'applique à des particules ponctuelles, c'est-à-dire dont l'extension spatiale est négligeable devant les distances entre particules.

Ce modèle des charges ponctuelles sera valable tant que les distances entre particules seront grandes par rapport aux dimensions de la charge. Dans ce cas, il est parfaitement justifié de considérer les particules chargées comme des points matériels portant une charge : leur extension spatiale est négligeable devant les autres distances.

Lorsqu'on observe la matière à l'échelle des atomes (par exemple avec un microscope à effet tunnel), on peut distinguer chaque atome. Il s'agit d'une observation à l'échelle microscopique. On caractérise cette échelle par une distance caractéristique d . Pour pouvoir conserver le caractère ponctuel des charges, cette distance doit être grande devant la taille a de l'atome. La distance d est typiquement de l'ordre de quelques dizaines de nanomètres pour les liquides et les solides et a de l'ordre de 0,1 nm donc $a \ll d$.

Lorsque l'observation est faite au niveau de l'expérience et que les atomes ne peuvent plus être distingués, on effectue une étude à l'échelle macroscopique. La distance caractéristique D de cette échelle est très grande devant d . Le fait qu'on ne peut plus distinguer les atomes incite à appliquer une description continue comme comportement moyen de la description microscopique locale.

Cependant cette approximation continue ne sera acceptable que s'il est possible de définir une échelle intermédiaire dite échelle mésoscopique. Elle est caractérisée par une distance l telle que :

- $d \ll l$ de sorte que la notion de moyenne ait un sens statistique,
- $l \ll D$ de façon à avoir une description suffisamment précise du milieu à l'aide de ces moyennes.

On obtient ainsi un lissage des grandeurs permettant une description continue. Dans la suite, on supposera que cette approximation continue est possible et que les conditions énoncées précédemment sont vérifiées.

1.2 Densité volumique de charges

On considère un volume élémentaire $d\tau_P$ autour d'un point P de charge totale dQ_P .

On définit la densité volumique de charges au point P , $\rho(P)$, par :

$$dQ_P = \rho(P)d\tau_P$$

La densité volumique de charges est, dans le cas général, une fonction du point P . Si elle est indépendante de P , on dit qu'elle est *uniforme*.

La charge totale s'obtient en sommant sur l'ensemble des volumes élémentaires les charges de chacun d'eux. Ceci se traduit mathématiquement par l'expression suivante :

$$Q = \iiint_{P \in \mathcal{V}} \rho(P) d\tau_P$$

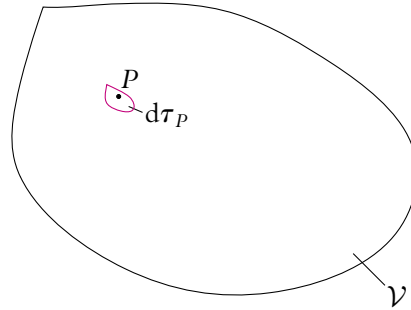


Figure 42.1 Volume élémentaire.

Exemple de densité volumique de charges

Pour une sphère de rayon R uniformément chargée en volume, on a :

$$\rho = \frac{Q}{\mathcal{V}} = \frac{3Q}{4\pi R^3}$$

Si on applique ce modèle au noyau d'hydrogène dont la charge vaut $Q = 1,6 \cdot 10^{-19}$ C et le rayon $R = 10^{-15}$ m (ou 1 fermi), on obtient :

$$\rho = 3,5 \cdot 10^{27} \text{ C.m}^{-3}$$

1.3 Densité surfacique de charges

Si on considère le cas d'une sphère conductrice, l'expérience montre¹ que les charges ne se répartissent pas dans la totalité du volume mais se localisent au voisinage de la surface de séparation entre la sphère et l'extérieur. On peut adopter le modèle suivant : la charge volumique ne dépend, compte tenu de la symétrie sphérique, que de la distance au centre de la sphère et suit une loi du type :

$$\rho(r) = \begin{cases} \rho_0 \exp\left(-\frac{R-r}{a}\right) & \text{si } 0 < r < R \\ 0 & \text{si } r > R \end{cases}$$

avec $a \ll R$. À une distance de quelques fois a du bord de la sphère, on aura $\rho(r) \simeq 0$ pour $0 < r < R$. Dans ce cas, il est tentant au niveau macroscopique de négliger l'épaisseur a de la répartition de charges. On est amené alors à considérer une distribution surfacique et non pas volumique.

¹ Le résultat sera établi dans le cours de deuxième année pour l'option MP/MP*.

La charge dQ_P peut s'écrire sous la forme :

$$dQ_P = \rho(P)d\tau_P = \rho(P)adS_P = \sigma(P)dS_P$$

en notant dS_P l'élément de surface élémentaire autour du point P . On définit ainsi la densité surfacique de charges

$$\sigma(P) = \lim_{a \rightarrow 0} (a\rho(P))$$

La charge totale s'obtient alors par l'intégrale :

$$Q = \iint_{P \in S} \sigma(P)dS_P$$

Ce type de modélisation conviendra bien pour décrire la distribution de charges au sein d'un conducteur. Son inconvénient est de créer une discontinuité artificielle de charges au niveau de la surface entre le conducteur et l'extérieur. Ceci peut se résoudre, si le besoin s'en fait sentir, en introduisant au niveau local une densité volumique de charges ρ sur une faible épaisseur a telle que $\rho = \frac{\sigma}{a}$. Il faudra alors avoir une estimation correcte de la valeur de a , ce qui justifiera ou non l'introduction de cette densité volumique de charges. On reviendra plus en détail sur ce type de problème dans le chapitre 45 sur un exemple et dans le cours de deuxième année.

1.4 Densité linéique de charges

De la même façon, les charges peuvent se localiser le long d'un fil dans le cas, par exemple, d'un conducteur très long pour lequel on peut négliger la section devant la longueur.

Ceci conduit à introduire des répartitions linéiques de charges :

$$dQ_P = \lambda(P)dl_P$$

en notant dl_P la longueur élémentaire le long du fil à proximité du point P et $\lambda(P)$ la densité linéique de charges.

La charge s'obtient par l'intégrale :

$$Q = \int_{P \in \mathcal{L}} \lambda(P)dl_P$$

Là encore, l'inconvénient est l'introduction de discontinuités qu'on peut résoudre comme précédemment en réintroduisant une densité volumique de charges au niveau local.

1.5 Charges ponctuelles

Dans certains cas, le volume de la distribution est très faible devant les dimensions caractéristiques du problème. Dans ce cas, on peut considérer que la charge est concentrée en un point et non répartie dans un volume. On arrive ainsi au modèle de la charge ponctuelle.

La charge totale s'obtient par une somme discrète sur les différentes charges ponctuelles :

$$Q = \sum_i q_i$$

L'inconvénient de ce type de distributions est encore l'introduction de discontinuités. En conclusion, toutes les distributions hormis la distribution volumique de charges introduisent des discontinuités. Ceci se justifie par le fait que les distributions sont naturellement volumiques et que l'utilisation d'autres types de distributions, pour les raisons évoquées précédemment, conduit à avoir en contrepartie des discontinuités. Celles-ci n'ont pas d'existence réelle et peuvent être résolues en réintroduisant les distributions volumiques de charges sous-jacentes au niveau local sous réserve de pouvoir estimer les dimensions caractéristiques de cette zone locale.

2. Déplacements, surfaces et volumes élémentaires

On vient d'introduire des densités volumiques, surfaciques et linéiques de charges ; il sera donc nécessaire d'exprimer dans les différents systèmes de coordonnées les déplacements, les surfaces et les volumes élémentaires.

Les déplacements élémentaires s'obtiennent en faisant varier de manière élémentaire chaque paramètre des coordonnées considérées. On les a déjà exprimés dans les chapitres concernant la cinématique de première période et les compléments de seconde période (dans le cours de mécanique). On se limitera ici à rappeler leurs expressions. En revanche, on explicitera les surfaces et les volumes élémentaires induits par ces déplacements.

2.1 Coordonnées cartésiennes

a) Déplacement élémentaire

Le déplacement élémentaire d'un point M de coordonnées (x, y, z) correspond à son déplacement jusqu'au point M' $(x + dx, y + dy, z + dz)$. On a donc :

$$d\overrightarrow{OM} = dx\overrightarrow{u_x} + dy\overrightarrow{u_y} + dz\overrightarrow{u_z}$$

b) Volume élémentaire

Comme le montre la figure ci-contre, les déplacements élémentaires suivant chaque axe induisent un volume élémentaire :

$$d\tau = dx dy dz$$

c) Surfaces élémentaires

De même, si on ne considère les déplacements que sur deux axes, on obtient trois surfaces élémentaires :

- $dx dy$, perpendiculaire au vecteur $\vec{u}_z$;
- $dx dz$, perpendiculaire au vecteur $\vec{u}_y$;
- $dy dz$, perpendiculaire au vecteur $\vec{u}_x$.

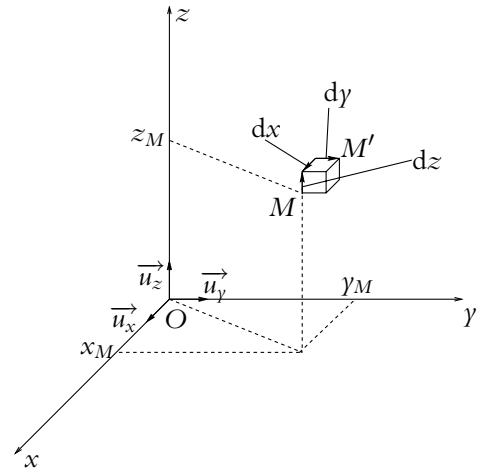


Figure 42.2 Volume et surfaces élémentaires en coordonnées cartésiennes.

2.2 Coordonnées cylindriques**a) Déplacement élémentaire**

Le déplacement élémentaire s'écrit :

$$d\vec{OM} = dr \vec{u}_r + r d\theta \vec{u}_\theta + dz \vec{u}_z$$

b) Volume élémentaire

Comme le montre la figure ci-contre, les déplacements élémentaires suivant chaque axe induisent un volume élémentaire :

$$d\tau = dr \times r d\theta \times dz$$

c) Surfaces élémentaires

De même, si on ne considère les déplacements que sur deux axes, on obtient trois surfaces élémentaires :

- $dr \times r d\theta$, perpendiculaire au vecteur $\vec{u}_z$;
- $dr \times dz$, perpendiculaire au vecteur $\vec{u}_\theta$;
- $r d\theta \times dz$, perpendiculaire au vecteur $\vec{u}_r$.

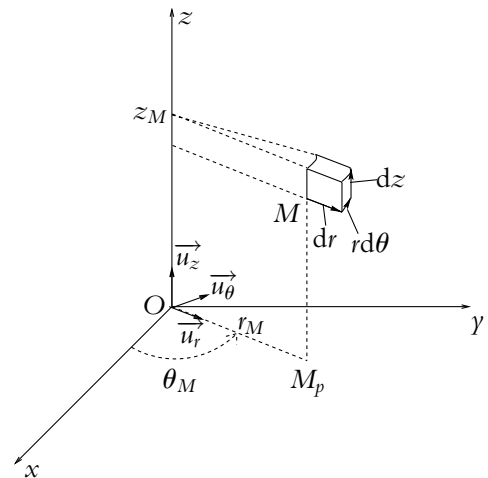


Figure 42.3 Volume et surfaces élémentaires en coordonnées cylindriques.

2.3 Coordonnées sphériques

a) Déplacement élémentaire

Le déplacement élémentaire s'écrit :

$$d\vec{OM} = dr\vec{u}_r + r d\theta\vec{u}_\theta + r \sin \theta d\varphi\vec{u}_\varphi$$

b) Volume élémentaire

Comme le montre la figure ci-contre, les déplacements élémentaires suivant chaque axe induisent un volume élémentaire :

$$d\tau = dr \times r d\theta \times r \sin \theta d\varphi$$

c) Surfaces élémentaires

De même, si on ne considère les déplacements que sur deux axes, on obtient trois surfaces élémentaires :

- $dr \times r d\theta$, perpendiculaire au vecteur $\vec{u}_\varphi$;
- $dr \times r \sin \theta d\varphi$, perpendiculaire au vecteur $\vec{u}_\theta$;
- $r d\theta \times r \sin \theta d\varphi$, perpendiculaire au vecteur $\vec{u}_r$.

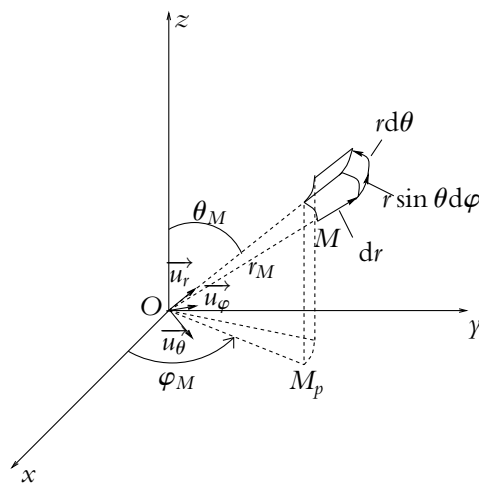


Figure 42.4 Volume et surfaces élémentaires en coordonnées sphériques.

3. Invariances par transformation géométrique

3.1 Invariance par translation

Soit une distribution de charges $\mathcal{D}$.

Elle sera dite invariante par translation si on retrouve la même distribution après avoir effectué la dite translation.

Pour une distribution d'extension finie, cela est exclu : aucune distribution ne peut vérifier cette propriété puisque seule une partie de la distribution tradatée correspond à la distribution initiale (avant translation).

Cependant, il est souvent utile d'envisager des distributions idéalisées pour simplifier l'analyse des problèmes. Par exemple, lorsque la distance du point considéré à la distribution est faible devant la taille de la distribution, on considérera que les dimensions de la distribution tendent vers l'infini. C'est la même chose pour un bateau perdu au milieu de l'océan : il ne voit pas les côtes et peut considérer l'océan comme un espace infini.

Premier exemple : « fil infini »

Soient un fil de longueur L et un point M distant de d de ce fil.

Si $d \ll L$, on peut considérer que le point M « voit » un fil de longueur infinie. On peut alors considérer l'invariance par translation le long de ce fil.

Cela conduira dans l'étude quantitative à choisir un système de coordonnées privilégiant un axe, à savoir les coordonnées cartésiennes ou cylindriques.

De plus, l'invariance par translation implique de retrouver la même situation avant et après translation. La variable d'espace décrivant cette translation n'est donc pas une variable pertinente dont dépendent les grandeurs étudiées. On n'aura plus de dépendance avec la composante le long de l'axe.

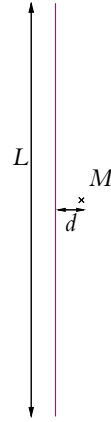


Figure 42.5 Fil infini.

Deuxième exemple : « plan infini »

Soit un plan de longueur L et de largeur l . On note M un point situé à une distance d de ce plan.

Si $d \ll L$ et $d \ll l$, on peut considérer que le point M « voit » un plan infini. On peut alors considérer les invariances par translation dans ce plan.

Pour l'étude quantitative, on choisira les coordonnées cartésiennes. En notant (xOy) le plan « infini », les invariances par translation parallèlement au plan auront pour effet de supprimer les dépendances en x et en y .

3.2 Invariance par rotation**a) Autour d'un axe**

Soit une distribution de charges $\mathcal{D}$ de densité volumique de charges $\rho(P)$.

Soit un point P quelconque de cette distribution $\mathcal{D}$. On note P' le point obtenu en effectuant une rotation d'angle θ autour d'un axe fixe Δ .

Si P' appartient à la distribution $\mathcal{D}$ et si la densité volumique de charges en P' est égale à celle en P :

$$\rho(P') = \rho(P)$$

alors la distribution $\mathcal{D}$ est invariante par la rotation d'angle θ autour de l'axe fixe Δ .

Exemple**Distribution constituée de quatre charges situées aux sommets d'un carré**

Si on considère l'axe Δ perpendiculaire au plan du carré et passant par son centre O , A a pour image B par la rotation d'angle $\frac{\pi}{2}$ et d'axe Δ , B pour image C , C pour image D et D pour image A . On a alors invariance de la distribution par rotation d'angle $\frac{\pi}{2}$ et d'axe Δ à condition que les quatre charges soient identiques pour vérifier l'égalité des charges avant et après rotation.

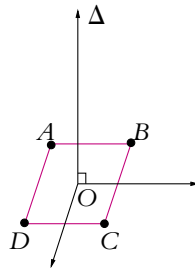


Figure 42.6 Distribution discrète de quatre charges ponctuelles.

Si cette invariance par rotation d'axe Δ et d'angle θ est vérifiée pour toute valeur de θ , on dit que la distribution est invariante par rotation autour de l'axe fixe Δ . C'est par exemple le cas d'un cercle uniformément chargé.

Dit autrement, l'image de la distribution $\mathcal{D}$ par n'importe quelle rotation autour de l'axe Δ est la distribution $\mathcal{D}$ elle-même.

Dans ce cas, on choisira les coordonnées cylindriques qui privilégient un axe ; cet axe sera celui autour duquel on a invariance par rotation. La composante angulaire notée généralement θ n'est plus alors un paramètre pertinent : les différentes quantités qu'on pourra être amené à calculer sont indépendantes de θ .

Par exemple, dans le cas du fil « infini », on a déjà obtenu l'indépendance vis-à-vis de la composante le long de l'axe (noté en général z) dans le paragraphe précédent. On obtient ici celle vis-à-vis de la composante angulaire θ . La seule variable pertinente est donc la distance r du point à l'axe.

b) Autour d'un point

Soit une distribution de charges $\mathcal{D}$ de densité volumique de charges $\rho(P)$.

Soit un point P quelconque de cette distribution $\mathcal{D}$. On note P' le point obtenu en effectuant une rotation d'angle θ autour d'un point fixe O .

Si P' appartient à la distribution $\mathcal{D}$ et si la densité volumique de charges en P' est égale à celle en P :

$$\rho(P') = \rho(P)$$

alors la distribution $\mathcal{D}$ est invariante par rotation d'angle θ autour du point fixe O .

Si l'invariance existe pour toute valeur de θ , on dit qu'on a invariance de la distribution par rotation autour du point O .

Dit autrement, l'image de la distribution $\mathcal{D}$ par une rotation quelconque autour du point O est la distribution $\mathcal{D}$ elle-même.

On se placera alors en coordonnées sphériques où aucun axe n'est privilégié. On aura alors l'indépendance vis-à-vis de toute composante angulaire (θ ou φ habituellement). Le seul paramètre sera la distance r du point au point O .

3.3 Remarque fondamentale

Il faut bien noter que les seuls paramètres géométriques des distributions ne suffisent pas pour avoir invariance. Il est nécessaire que l'image d'un point de la distribution soit un autre point de la distribution après transformation mais cela n'est pas suffisant. Il faut aussi qu'au point obtenu après transformation, on ait la même charge ou la même densité de charges qu'au point initial.

Il y a donc deux conditions à vérifier :

- l'image par une transformation (translation ou rotation) d'un point de la distribution doit appartenir à la distribution,
- les charges ou densités de charges au point et en son image par la transformation doivent être les mêmes.

Dans le cas des distributions uniformément réparties, il n'y a aucun problème mais dans le cas de distributions quelconques, il faudra être bien sûr de cet aspect pour utiliser les propriétés d'invariance.

3.4 Invariances et choix des coordonnées

Les invariances des distributions de charges (ou sources) impliquent un choix adapté des coordonnées à employer :

- en cas d'invariance par translation, utilisation de coordonnées privilégiant un axe donc des coordonnées cartésiennes ou cylindriques,
- en cas d'invariance par rotation, utilisation de coordonnées précisant un angle donc des coordonnées cylindriques ou sphériques.

43

Champ électrostatique

Ce chapitre énonce la loi phénoménologique de Coulomb donnant la force exercée par une particule chargée sur une autre. Cela permet d'introduire la notion de champ électrostatique. L'étude des symétrie des distributions de charges permettra de simplifier efficacement celle du champ électostatique.

1. Loi de Coulomb

Il s'agit d'une loi phénoménologique : on constate expérimentalement que les actions entre corps électrisés obéissent à ce modèle mais cela n'a pas été démontré à partir d'autres principes.

1.1 Influence de la distance

Charles-Augustin de Coulomb a établi expérimentalement que l'interaction électrostatique est une interaction newtonienne, c'est-à-dire une interaction à force centrale qui décroît en $\frac{1}{r^2}$ où r est la distance entre les deux charges électriques en interaction.

1.2 Principe des actions réciproques

Les charges sont supposées ponctuelles : le principe de l'action et de la réaction ou des actions réciproques s'applique. Si une charge q_1 exerce une force $\vec{F}_{1 \rightarrow 2}$ sur la charge q_2 , q_2 exerce la force $\vec{F}_{2 \rightarrow 1}$ vérifiant :

$$\vec{F}_{2 \rightarrow 1} = -\vec{F}_{1 \rightarrow 2}$$

La force exercée par la charge q_1 est proportionnelle à q_1 et la force exercée par la charge q_2 proportionnelle à q_2 . Ceci se traduit par : $F_{1 \rightarrow 2} = \alpha q_1$ et $F_{2 \rightarrow 1} = \alpha' q_2$. Comme le principe des actions réciproques s'écrit : $F_{1 \rightarrow 2} = F_{2 \rightarrow 1}$, on en déduit

$\alpha q_1 = \alpha' q_2$. La constante α est donc proportionnelle à q_2 et α' à q_1 : on en déduit que $F_{1 \rightarrow 2}$ et $F_{2 \rightarrow 1}$ sont proportionnelles au produit $q_1 q_2$.

D'autre part, le principe de l'action et de la réaction impose que la force d'interaction soit portée par la droite $M_1 M_2$.

1.3 Énoncé de la loi de Coulomb

La dépendance en $\frac{1}{r^2}$ de la force de Coulomb et le principe des actions réciproques permettent de justifier l'allure de la loi de Coulomb.

Soient deux charges ponctuelles q_1 et q_2 situées respectivement en M_1 et M_2 . La charge q_1 exerce sur la charge q_2 la force

$$\vec{F}_{1 \rightarrow 2} = \frac{1}{4\pi\epsilon_0} \frac{q_1 q_2}{(M_1 M_2)^3} \overrightarrow{M_1 M_2}$$

On peut noter que le coefficient de proportionnalité $\frac{1}{4\pi\epsilon_0}$ fixe l'unité de la charge q , dans les unités du Système International :

$$\epsilon_0 = 8,85 \cdot 10^{-12} \text{ C}^2 \cdot \text{s}^2 \cdot \text{kg}^{-1} \cdot \text{m}^{-3} \quad (\text{ou } \text{F} \cdot \text{m}^{-1})$$

ou :

$$\frac{1}{4\pi\epsilon_0} = 9 \cdot 10^9 \text{ F}^{-1} \cdot \text{m}$$

ϵ_0 est appelé *permittivité diélectrique du vide*.

Cette loi vérifie les propriétés établies précédemment à partir des observations.

1.4 Principe de superposition

Soit une charge q placée en M .

Dans un premier temps, on place en M_1 une charge q_1 . La charge q subit alors la force :

$$\vec{F}_1 = \frac{1}{4\pi\epsilon_0} \frac{q q_1}{(M M_1)^3} \overrightarrow{M_1 M}$$

Dans une autre expérience, on place en M_2 une charge q_2 . La charge q subit alors la force :

$$\vec{F}_2 = \frac{1}{4\pi\epsilon_0} \frac{q q_2}{(M M_2)^3} \overrightarrow{M_2 M}$$

Si on place maintenant simultanément une charge q_1 en M_1 et une charge q_2 en M_2 , la présence de q_2 ne modifie pas l'action de q_1 et réciproquement. On peut alors, du fait de l'addition vectorielle des forces (contenue dans le principe fondamental de la

dynamique), déduire la force $\vec{F}$ subie par la charge q :

$$\vec{F} = \vec{F}_1 + \vec{F}_2$$

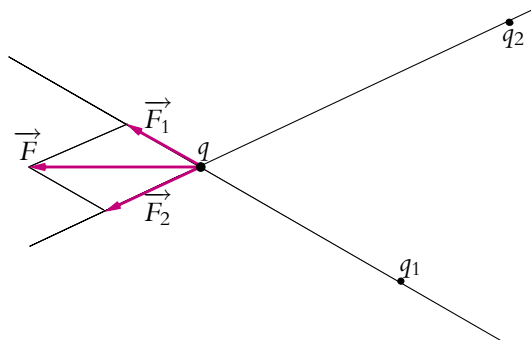


Figure 43.1 Principe de superposition avec des charges de même nature.

L'étude de l'interaction entre plusieurs charges est alors considérablement simplifiée : il suffit de se ramener au problème de l'interaction entre deux charges ponctuelles et de sommer les effets. Le point important est l'indépendance de l'interaction de deux particules chargées vis-à-vis de la présence ou non d'autres charges.

1.5 Définition expérimentale de la charge

Pour un corps macroscopique, il est impensable de connaître avec précision la valeur de sa charge totale Q à partir de l'ensemble des charges qui le constituent. Ces dernières sont en trop grand nombre (de l'ordre du nombre d'Avogadro). La définition pratique de la charge se fait donc par comparaison des charges entre elles. Pour cela, on soumet une charge q à l'action d'une autre charge q_1 : elle est alors soumise à la force $\vec{F}_1$. Soumise à l'action d'une charge q_2 , elle subit la force $\vec{F}_2$.

On en déduit en appliquant la loi de Coulomb :

$$\frac{q_2}{q_1} = \frac{F_2}{F_1}$$

Connaissant la valeur de q_1 et les modules des forces $\vec{F}_1$ et $\vec{F}_2$, on en déduit la valeur de q_2 . Cela permet de définir expérimentalement la charge.

2. Notion de champ

2.1 Définition

On désigne par champ scalaire (respectivement vectoriel) une grandeur scalaire (respectivement vectorielle) définie en chaque point de l'espace ou d'un domaine de l'espace. Il faudra, outre l'expression de ce champ, également préciser l'espace sur

lequel il est défini ainsi que les valeurs prises par ce dernier aux limites du domaine : on parle de conditions aux limites.

Par exemple, on peut définir un champ de pesanteur, un champ de température, etc.

2.2 Champ uniforme

Un champ sera dit *uniforme* si la grandeur le définissant prend la même valeur en tout point du domaine sur lequel on le définit.

2.3 Champ stationnaire

Un champ sera dit *stationnaire*¹ si la grandeur le définissant garde la même valeur au cours du temps, autrement dit si la grandeur ne dépend pas du temps.

3. Définition du champ électrostatique

3.1 Champ électrostatique créé par une charge ponctuelle

La loi de Coulomb s'écrit :

$$\vec{F}_{1 \rightarrow 2} = \frac{1}{4\pi\epsilon_0} \frac{q_1 q_2}{(M_1 M_2)^3} \overrightarrow{M_1 M_2}$$

qu'on peut mettre sous la forme :

$$\vec{F}_{1 \rightarrow 2} = q_2 \left(\frac{1}{4\pi\epsilon_0} \frac{q_1}{(M_1 M_2)^3} \overrightarrow{M_1 M_2} \right) = q_2 \vec{E}_1(M_2)$$

en notant :

$$\vec{E}_1(M_2) = \frac{1}{4\pi\epsilon_0} \frac{q_1}{(M_1 M_2)^3} \overrightarrow{M_1 M_2} \quad (43.1)$$

le champ électrostatique créé au point M_2 par la charge électrique q_1 située en M_1 .

On constate que, par définition, ce champ ne dépend pas de la charge q_2 mais uniquement de l'endroit où celle-ci se trouve. On peut donc parfaitement définir un champ électrostatique indépendamment des charges sur lesquelles il s'applique. L'important est l'effet électrostatique créé par une distribution.

On notera qu'il est possible d'écrire ce champ sous la forme :

$$\vec{E}_1(M_2) = \frac{1}{4\pi\epsilon_0} \frac{q_1}{r^2} \vec{u}_r$$

où r désigne la distance entre M_1 et M_2 et $\vec{u}_r$ le vecteur unitaire directeur de la droite $M_1 M_2$ orienté dans le sens de M_1 vers M_2 . On fait bien apparaître ainsi la

¹ On dit parfois aussi *permanent*.

dépendance en $\frac{1}{r^2}$. On évitera cependant de retenir l'expression du champ sous cette forme : son orientation n'est pas suffisamment explicite. Il est préférable de retenir l'expression (43.1).

3.2 Application du principe de superposition

Soit une charge q placée en un point M de l'espace.

Soit un ensemble de charges ponctuelles q_i placées aux points P_i . En appliquant le principe de superposition, la force exercée par l'ensemble des charges q_i sur la charge q s'écrit :

$$\vec{F} = \sum_i \frac{1}{4\pi\epsilon_0} \frac{qq_i}{(P_iM)^3} \overrightarrow{P_iM}$$

qu'on peut écrire :

$$\vec{F} = q \sum_i \frac{1}{4\pi\epsilon_0} \frac{q_i}{(P_iM)^3} \overrightarrow{P_iM} = q \vec{E}(M)$$

L'application du principe de superposition conduit à dire que le champ électrostatique créé par une distribution de charges ponctuelles q_i placées en des points P_i est égal à la somme des champs électrostatiques créés par chacune des charges q_i considérée seule. Le principe de superposition s'applique donc également au champ électrostatique.

$$\vec{E}(M) = \frac{1}{4\pi\epsilon_0} \sum_i \frac{q_i}{(P_iM)^3} \overrightarrow{P_iM}$$

3.3 Champ créé par une distribution quelconque de charges

Les définitions précédentes du champ électrostatique ne concernent que les distributions discrètes de charges. On peut généraliser ces résultats aux distributions continues (volumique, surfacique ou linéique).

a) Distribution volumique de charges

À partir de la densité volumique de charges $\rho(P)$ qui dépend du point de l'espace P où elle se trouve, on obtient, en sommant les différentes contributions, l'expression du champ électrostatique :

$$\vec{E}(M) = \frac{1}{4\pi\epsilon_0} \iiint_{P \in \mathcal{V}} \rho(P) \frac{\overrightarrow{PM}}{(PM)^3} d\tau_P$$

Il faut noter que l'intégration porte sur le volume où sont situées les charges créant le champ électro-

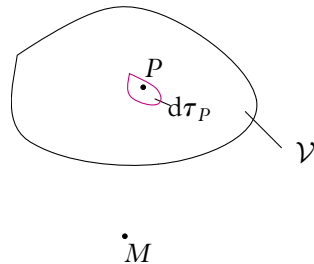


Figure 43.2 Distribution volumique de charges.

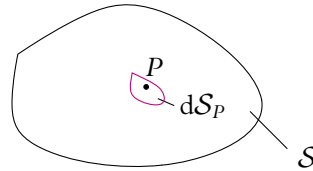
statique et non sur l'endroit où on cherche l'expression du champ. Le point M est le point où on cherche le champ ; il peut appartenir à la distribution ou non. Le champ électrostatique dépend du point M où on le calcule : on le rappelle par la notation $\vec{E}(M)$. Le point P décrit toute la distribution de charges $\mathcal{V}$. La densité volumique de charges $\rho(P)$ dépend du point P . Il est donc fortement conseillé de ne pas omettre M et P quand on apprend l'expression du champ.

Il faut également faire attention au fait que l'intégration concerne des vecteurs : le vecteur $\vec{PM}$ varie avec le point P sur lequel porte l'intégration. Il est donc nécessaire de projeter l'expression à intégrer sur une base fixe et de calculer par intégration chacune des composantes non nulles².

b) Distribution surfacique de charges

En utilisant les mêmes notations que celles employées auparavant, on obtient :

$$\vec{E}(M) = \frac{1}{4\pi\epsilon_0} \iint_{P \in S} \sigma(P) \frac{\vec{PM}}{(PM)^3} dS_P$$



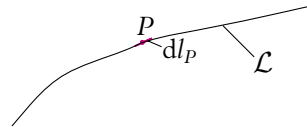
•M

Figure 43.3 Distribution surfacique de charges.

c) Distribution linéique de charges

De même, pour une distribution linéique de charges $\lambda(P)$:

$$\vec{E}(M) = \frac{1}{4\pi\epsilon_0} \int_{P \in \mathcal{L}} \lambda(P) \frac{\vec{PM}}{(PM)^3} dl_P$$



•M

Figure 43.4 Distribution linéique de charges.

Toutes ces expressions, que les distributions soient volumiques, surfaciques ou linéiques, ne sont valables que pour des distributions d'extension finie dans l'espace.

4. Analogie avec le champ de gravitation

4.1 Rappel : force de gravitation

Lors de la présentation des interactions, on a déjà parlé de la force de gravitation qui s'exerce entre deux points matériels M_1 et M_2 de masse respective m_1 et m_2 .

² Les composantes non nulles auront été déterminées auparavant lors de l'analyse des symétries.

L'expression de la force exercée par M_1 sur M_2 est :

$$\vec{F}_{1 \rightarrow 2} = -G \frac{m_1 m_2}{(M_1 M_2)^3} \overrightarrow{M_1 M_2}$$

en notant $G = 6,67 \cdot 10^{-11} \text{ N.m}^2.\text{kg}^{-2}$ la constante de gravitation universelle.

Cette force est formellement analogue à la force électrostatique. La seule différence provient du fait qu'une masse est toujours positive tandis qu'une charge peut être négative ou positive : la force gravitationnelle est toujours attractive tandis que la force électrostatique est attractive lorsque les charges sont de signe opposé et répulsive sinon.

Le tableau suivant explicite l'analogie :

gravitation	électrostatique
masse m	charge électrique q
force $\vec{F}_{1 \rightarrow 2} = -G \frac{m_1 m_2}{(M_1 M_2)^3} \overrightarrow{M_1 M_2}$	force $\vec{F}_{1 \rightarrow 2} = \frac{1}{4\pi\epsilon_0} \frac{q_1 q_2}{(M_1 M_2)^3} \overrightarrow{M_1 M_2}$
constante universelle de gravitation : $-G = -6,67 \cdot 10^{-11} \text{ N.m}^2.\text{kg}^{-2}$	constante dépendant de la permittivité dans le vide ϵ_0 : $\epsilon_0 = 8,85 \cdot 10^{-12} \text{ F.m}^{-1}$

4.2 Champ de gravitation

a) Définition du champ gravitationnel créé par une masse ponctuelle

Comme pour la force électrostatique, on peut donc introduire la notion de champ. Le champ de gravitation créé en M_2 par une masse ponctuelle m_1 située en M_1 s'écrit donc :

$$\vec{A}_1(M_2) = -G \frac{m_1}{(M_1 M_2)^3} \overrightarrow{M_1 M_2}$$

Comme le champ électrostatique, ce champ ne dépend pas de la masse m_2 mais uniquement de l'endroit où elle se trouve.

b) Distributions de masse

Comme pour les charges, on définit des densités volumiques, surfaciques ou linéiques de masse :

- distribution volumique de masses : dans un volume élémentaire $d\tau_P$ autour du point P , on a une masse dm_P et on définit la densité volumique de masse ou *masse volumique* $\rho(P)$ par :

$$dm_P = \rho(P)d\tau_P$$

- distribution surfacique de masses : lorsque l'une des dimensions de la distribution de masse est négligeable devant les autres, on peut considérer une distribution surfacique ; pour une surface élémentaire dS_P autour du point P de masse dm_P , on définit la densité surfacique de masse ou *masse surfacique* $\sigma(P)$ par :

$$dm_P = \sigma(P)dS_P$$

- distribution linéique de masses : lorsque la masse se répartit le long d'un fil de section négligeable, on peut considérer une distribution linéique ; pour une longueur élémentaire dl_P à proximité du point P de masse dm_P , on définit la densité linéique de masse ou *masse linéique* $\lambda(P)$ par :

$$dm_P = \lambda(P)dl_P$$

Les masses totales s'obtiennent en sommant les différents éléments, ce qui se traduit mathématiquement par une intégrale.

On peut enfin considérer que les masses sont ponctuelles lorsque leur volume est très faible devant les dimensions caractéristiques du problème. On a alors une distribution discrète.

c) Champ gravitationnel créé par une distribution quelconque de masses

En utilisant le principe de superposition comme pour le champ électrostatique, on obtient les expressions du champ gravitationnel créé par n'importe quelle distribution de masses :

- distribution volumique de masses :

$$\vec{A}(M) = -G \iiint_{P \in \mathcal{V}} \rho(P) \frac{\vec{PM}}{(PM)^3} d\tau_P$$

- distribution surfacique de masses :

$$\vec{A}(M) = -G \iint_{P \in \mathcal{S}} \sigma(P) \frac{\vec{PM}}{(PM)^3} dS_P$$

- distribution linéique de masses :

$$\vec{A}(M) = -G \int_{P \in \mathcal{L}} \lambda(P) \frac{\vec{PM}}{(PM)^3} dl_P$$

- distribution discrète de masses :

$$\vec{A}(M) = -G \sum_i m_i \frac{\vec{P_i M}}{(P_i M)^3}$$

5. Invariances, symétries et propriétés du champ

5.1 Principe de Curie

L'importance de la recherche des invariances et des symétries des distributions est liée au principe de Curie : les effets et les causes présentent les mêmes invariances et les mêmes symétries. On peut l'énoncer de la façon suivante :

Les éléments d'invariance et de symétrie des causes doivent se retrouver dans les effets produits.

Ici les effets sont le champ électrique $\vec{E}$ et le potentiel électrostatique V (qui sera défini au chapitre 44) et les causes, les sources du champ et du potentiel, à savoir les distributions de charges. On va donc chercher les invariances et les symétries des distributions de charges, ce qui permettra de connaître celles du champ et du potentiel électrique.

5.2 Conséquences des invariances de la distribution

a) Invariance par translation

Si la distribution de charges est invariante par toute translation le long de l'axe Ox alors, par application du principe de Curie, on a :

$$\forall x, \forall a \quad \vec{E}(x, y, z) = \vec{E}(x + a, y, z)$$

x n'est plus une variable pertinente et on en déduit que le champ électrostatique ne dépend pas de cette variable :

$$\vec{E}(M) = \vec{E}(y, z)$$

b) Invariance par rotation

Si la distribution de charges est invariante par toute rotation d'angle θ (que ce soit par rapport à un axe ou à un point), on a :

$$\forall \theta, \forall \alpha \quad \left\| \vec{E}(\theta, \bullet) \right\| = \left\| \vec{E}(\theta + \alpha, \bullet) \right\|$$

en notant $\bullet$ toute autre variable dont on ne se préoccupe pas ici.

On en déduit que le champ ne dépend pas de θ .

L'invariance permet alors de supprimer une ou plusieurs coordonnées, ce qui réduit le nombre de variables d'espace à considérer.

5.3 Invariance des lois de l'électrostatique

Pour l'interaction électrostatique, le référentiel du laboratoire peut être considéré comme galiléen. En l'absence d'interaction, un corps animé d'une vitesse la conserve quelle que soit sa direction. L'espace est donc homogène isotrope : aucune direction et aucune origine n'est privilégiée.

De plus, la charge électrostatique est invariante par tout changement de référentiel et par tout déplacement.

Par conséquent, les lois de l'électrostatique doivent également respecter cette propriété d'invariance : quel que soit le déplacement envisagé, les lois de l'électrostatique sont les mêmes avant et après déplacement du point considéré.

Soit une charge ponctuelle q placée en P . Elle crée en un point M le champ électrostatique $\vec{E}(M)$:

$$\vec{E}(M) = \frac{1}{4\pi\epsilon_0} \frac{q}{(PM)^3} \overrightarrow{PM}$$

On fait subir une transformation au système que constituent la charge q et le point M où on cherche les valeurs du champ électrostatique. La charge (qui ne change pas de valeur du fait de son caractère invariant) se retrouve en P' et crée en M' , image de M par la transformation envisagée, un champ électrostatique $\vec{E}'(M')$:

$$\vec{E}'(M') = \frac{1}{4\pi\epsilon_0} \frac{q}{(P'M')^3} \overrightarrow{P'M'}$$

L'invariance des lois de l'électrostatique permet d'avoir la même écriture pour le champ électrostatique.

Dans le cas où la transformation est une isométrie (ce qui sera le cas pour les situations envisagées de translation, de rotation et de symétrie plane), on a :

$$PM = P'M'$$

d'où :

$$\vec{E}'(M') = \frac{1}{4\pi\epsilon_0} \frac{q}{(PM)^3} \overrightarrow{P'M'}$$

Le champ électrostatique se transforme comme le vecteur $\overrightarrow{PM}$. Le seul cas où on aura $\vec{E}'(M') = \vec{E}(M)$ est le cas d'une translation dont une des propriétés est de transformer un vecteur $\vec{u}$ en un vecteur $\vec{u}' = \vec{u}$.

5.4 Symétries des sources par rapport à un plan

Soit une charge q située en P et un point M où on cherche le champ électrostatique.

Par la symétrie par rapport au plan $\mathcal{P}$, la charge q devient la charge q' située en P' , symétrique de P par rapport au plan $\mathcal{P}$ et le point M devient le point M' , symétrique de M par rapport au plan $\mathcal{P}$.

On peut remarquer qu'une symétrie par rapport à un plan transforme une base directe en une base indirecte.

L'invariance de la charge s'écrit :

$$q = q'$$

L'invariance des lois de l'électrostatique permet alors d'écrire :

$$\vec{E}'(M') = \frac{1}{4\pi\epsilon_0} \frac{q}{(P'M')^3} \overrightarrow{P'M'}$$

La symétrie par rapport à un plan est une isométrie donc : $P'M' = PM$ et :

$$\vec{E}'(M') = \frac{1}{4\pi\epsilon_0} \frac{q}{(PM)^3} \overrightarrow{P'M'}$$

soit :

$$\vec{E}'(M') = \text{sym}_{\mathcal{P}} \left(\vec{E}(M) \right)$$

où $\text{sym}_{\mathcal{P}} \vec{E}(M)$ désigne le symétrique du vecteur $\vec{E}(M)$ par rapport au plan $\mathcal{P}$.

On généralise le résultat à une distribution quelconque grâce au principe de superposition. Le champ électrostatique créé par la distribution symétrique par rapport à un plan $\mathcal{P}$ de la distribution d'origine en un point M' symétrique d'un point M est le symétrique par rapport au plan $\mathcal{P}$ du champ électrostatique créé par la distribution d'origine au point M :

$$\vec{E}'(M') = \text{sym}_{\mathcal{P}} \left(\vec{E}(M) \right)$$

On est maintenant à même d'en tirer les conséquences pour des distributions présentant des plans de symétrie ou d'antisymétrie.

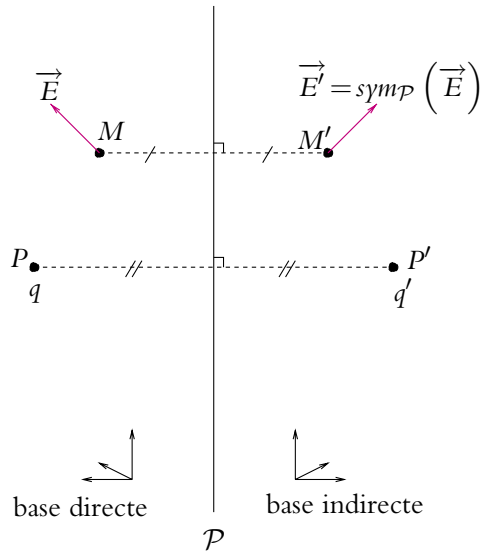


Figure 43.5 Champ électrostatique et symétrie par rapport à un plan.

5.5 Plans de symétrie et d'antisymétrie des sources

a) Plan de symétrie

On appelle *plan de symétrie* un plan tel que la distribution obtenue en déplaçant les charges selon une symétrie par rapport à ce plan est identique à la distribution initiale.

Exemple

Soit la distribution discrète composée de deux charges ponctuelles q_1 et q_2 représentée ci-contre.

Par symétrie par rapport au plan $\mathcal{P}$, la charge q_1 devient q_2 et la charge q_2 devient q_1 .

Si $q_1 = q_2$, le fait d'effectuer la symétrie par rapport au plan $\mathcal{P}$ ne modifie pas la distribution dans sa globalité. On retrouve la distribution initiale : la distribution possède le plan $\mathcal{P}$ pour plan de symétrie.

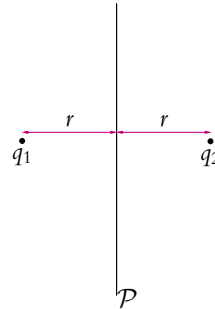


Figure 43.6 Plan de symétrie.

b) Plan d'antisymétrie

On appelle *plan d'antisymétrie* un plan tel que la distribution obtenue en déplaçant les charges selon une symétrie par rapport à ce plan est opposée à la distribution initiale.

Dans l'exemple précédent, si $q_1 = -q_2$, $\mathcal{P}$ est un plan d'antisymétrie.

Ces symétries permettront d'obtenir des informations supplémentaires intéressantes sur la direction des champs envisagés.

5.6 Distribution ayant un plan de symétrie

Soit une distribution $\mathcal{D}$ admettant le plan $\mathcal{P}$ comme plan de symétrie.

On vient d'établir que :

$$\vec{E}'(M') = \text{sym}_{\mathcal{P}} \left(\vec{E}(M) \right) \quad \text{où} \quad M' = \text{sym}_{\mathcal{P}}(M) \quad (43.2)$$

en notant $\vec{E}'$ le champ créé par la distribution $\mathcal{D}'$ symétrique de $\mathcal{D}$ par rapport au plan $\mathcal{P}$ et $\vec{E}$ le champ créé par la distribution $\mathcal{D}$.

Or, d'après le principe de Curie, le champ électrique créé en M' l'est soit par $\mathcal{D}$ soit par $\mathcal{D}'$, ce qui se traduit par

$$\vec{E}'(M') = \vec{E}(M')$$

puisque les densités de charges ρ' de $\mathcal{D}'$ et ρ de $\mathcal{D}$ vérifient $\rho' = \rho$ du fait que $\mathcal{P}$ est un plan de symétrie de $\mathcal{D}$.

En utilisant la relation (43.2), on obtient donc :

$$\vec{E}(M') = \text{sym}_{\mathcal{P}} \left(\vec{E}(M) \right)$$

Si M appartient à un plan de symétrie des sources, alors $M' = M$. On en déduit donc que :

$$\vec{E}(M) = \text{sym}_{\mathcal{P}} \left(\vec{E}(M) \right)$$

ce qui se traduit en utilisant les indices t et n respectivement pour les composantes tangentielle et normale au plan $\mathcal{P}$:

$$\begin{cases} \vec{E}_t(M) = \vec{E}_t(M) \\ \vec{E}_n(M) = -\vec{E}_n(M) \end{cases}$$

d'où

$$\vec{E}_n(M) = \vec{0}$$

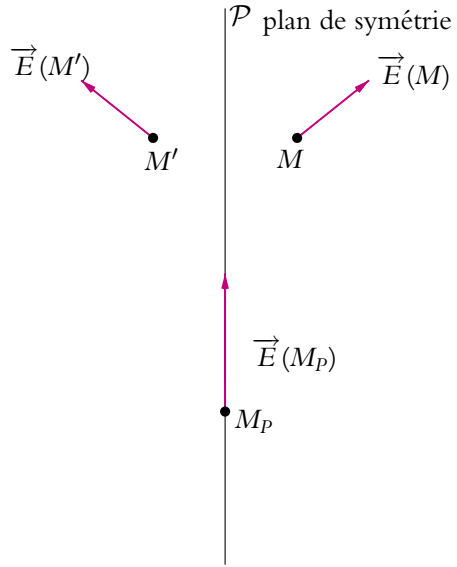


Figure 43.7 Champ électrostatique et plan de symétrie.

Le champ électrostatique n'a donc que des composantes tangentielles sur les plans de symétrie de la distribution qui le crée : il appartient aux plans de symétrie de la distribution de charges qui le crée.

5.7 Distribution ayant un plan d'antisymétrie

Soit une distribution $\mathcal{D}$ admettant le plan $\mathcal{P}$ comme plan d'antisymétrie.

On a toujours l'égalité (43.2).

Or d'après le principe de Curie, le champ électrique créé en M' l'est soit par $\mathcal{D}$ soit par $\mathcal{D}'$, ce qui se traduit par

$$\vec{E}'(M') = -\vec{E}(M')$$

puisque les densités de charges ρ' de $\mathcal{D}'$ et ρ de $\mathcal{D}$ vérifient $\rho' = -\rho$ du fait que $\mathcal{P}$ est un plan d'antisymétrie de $\mathcal{D}$.

En utilisant la relation (43.2), on obtient donc :

$$\vec{E}(M') = -\text{sym}_{\mathcal{P}} \left(\vec{E}(M) \right)$$

Si M appartient à un plan d'antisymétrie, $M' = M$.

On en déduit donc que :

$$\vec{E}(M) = -\text{sym}_{\mathcal{P}}(\vec{E}(M))$$

ce qui se traduit en utilisant les indices t et n respectivement pour les composantes tangentielle et normale au plan $\mathcal{P}$:

$$\begin{cases} \vec{E}_t(M) = -\vec{E}_t(M) \\ \vec{E}_n(M) = \vec{E}_n(M) \end{cases}$$

On en déduit :

$$\vec{E}_t(M) = \vec{0}$$

Le champ électrostatique n'a donc qu'une composante normale sur les plans d'antisymétrie de la distribution qui le crée : il est perpendiculaire aux plans d'antisymétrie de la distribution de charges qui le crée.

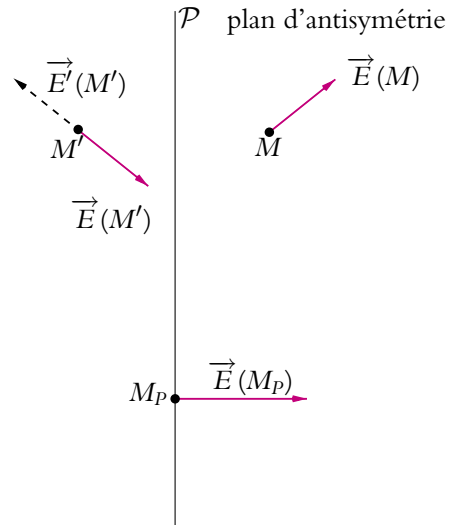


Figure 43.8 Champ électrostatique et plan d'antisymétrie.

5.8 Caractère polaire du champ électrostatique

L'analyse des plans de symétrie et d'antisymétrie permet donc de donner la direction du champ électrostatique :

- le champ électrostatique appartient aux plans de symétrie de la distribution de charges qui le crée,
- le champ électrostatique est perpendiculaire aux plans d'antisymétrie de la distribution de charges qui le crée.

Le champ électrostatique se transforme donc comme un vrai vecteur lors d'une symétrie par rapport à un plan : on a les mêmes lois de transformations que pour un vecteur « géométrique ». On dit que le champ électrostatique est un vecteur *polaire*.

Cette propriété peut paraître évidente, mais elle prendra tout son sens quand on étudiera le champ magnétique et qu'on verra que ce n'est pas un vecteur polaire mais un vecteur axial ou pseudo-vecteur.

5.9 Application au champ gravitationnel

Comme il n'existe pas de masse négative, seuls peuvent exister des plans de symétrie pour des distributions de masses.

En revanche, les résultats obtenus pour le champ électrostatique pour les distributions possédant un plan de symétrie s'appliquent au champ gravitationnel : le champ gravi-

tationnel appartient au plan de symétrie de la distribution de masses qui le crée. De part et d'autre d'un plan de symétrie de la distribution de masses qui le crée, le champ gravitationnel est symétrique soit :

$$\vec{A}(M') = \text{sym}_{\mathcal{P}} \left(\vec{A}(M) \right)$$

en notant M' le symétrique par rapport au plan $\mathcal{P}$ du point M et en supposant que la distribution $\mathcal{D}$ est invariante par la symétrie par rapport à ce plan.

6. Application au champ électrostatique créé par un segment dans son plan médiateur ou par un fil infini uniformément chargé

6.1 Champ créé par un segment uniformément chargé dans son plan médiateur

Soit un segment AB de longueur l uniformément chargé avec une densité linéique λ placé le long de l'axe Ox entre les abscisses $-\frac{l}{2}$ et $\frac{l}{2}$. On cherche le champ électrostatique créé par ce segment dans le plan médiateur yOz du segment AB .

a) Analyse des invariances

- Du fait de la symétrie de la distribution de charges, on utilise les coordonnées cartésiennes.
- On a invariance par rotation autour de l'axe Ox : tous les plans contenant l'axe Ox sont donc équivalents et s'obtiennent les uns à partir des autres par rotation autour de l'axe Ox . Par conséquent, on peut limiter l'étude à un plan contenant l'axe Ox , par exemple le plan xOy .
- On cherche le champ dans le plan médiateur donc le long de l'axe Oy d'après la remarque précédente. Le point M a pour coordonnées $(0, y, 0)$ donc

$$\vec{E}(M) = \vec{E}(y)$$

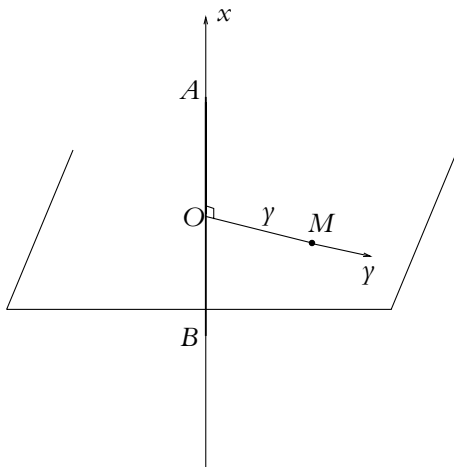


Figure 43.9

b) Analyse des symétries

Le plan médiateur du segment AB et le plan passant par M et contenant l'axe Ox sont des plans de symétrie donc le champ appartient à l'intersection de ces deux plans. D'après les remarques du paragraphe précédent, le champ est suivant Oy :

$$\vec{E}(M) = E(y)\vec{u}_y$$

c) Calcul direct du champ

L'expression générale du champ est :

$$\vec{E}(M) = \frac{1}{4\pi\epsilon_0} \int_{-\frac{l}{2}}^{\frac{l}{2}} \frac{\lambda dx}{(PM)^3} \vec{PM}$$

avec $PM = \sqrt{x^2 + y^2}$.

Du fait des résultats précédents, on ne calcule que la composante du champ suivant Oy :

$$E_y = \frac{1}{4\pi\epsilon_0} \int_{-\frac{l}{2}}^{\frac{l}{2}} \frac{\lambda \cos \alpha}{(PM)^2} dx$$

On préfère paramétrer le problème par l'angle α plutôt que par l'abscisse x du point P , ce qui rend les calculs plus simples. Il faut exprimer PM et dx en fonction de α et de y .

D'après la figure ci-contre, $PM = \frac{y}{\cos \alpha}$,

$\tan \alpha = \frac{x}{y}$ donc $\frac{d\alpha}{\cos^2 \alpha} = \frac{dx}{y}$. On en déduit :

$$E_y = \frac{1}{4\pi\epsilon_0} \int_{-\alpha_0}^{\alpha_0} \frac{\lambda \cos \alpha}{y} d\alpha$$

et

$$E_y = \frac{1}{2\pi\epsilon_0} \frac{\lambda \sin \alpha_0}{y}$$

$$\text{Or } \sin \alpha_0 = \frac{\frac{l}{2}}{\sqrt{\left(\frac{l}{2}\right)^2 + y^2}} = \frac{1}{\sqrt{1 + \frac{4y^2}{l^2}}},$$

on en déduit :

$$E_y = \frac{\lambda}{2\pi\epsilon_0} \frac{l}{y\sqrt{l^2 + 4y^2}}$$

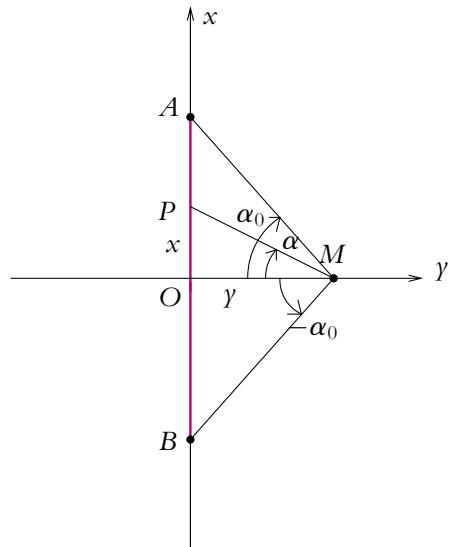


Figure 43.10

6.2 Champ électrostatique créé par un fil infini uniformément chargé

On calcule le champ en tout point M de l'espace.

a) Invariances et symétries

Du fait de la symétrie de la distribution de charges, on préfère utiliser les coordonnées cylindriques d'axe Ox .

L'étude des invariances et des symétries effectuée au paragraphe précédent pour un segment se généralise :

- l'invariance par rotation autour de l'axe Ox implique que la norme du champ ne dépend pas de l'angle polaire θ ;
- l'invariance par translation le long de l'axe Ox implique que la norme du champ ne dépend pas de l'abscisse x ;
- les deux remarques précédentes imposent que la norme du champ ne dépende que de r : $\|\vec{E}(M)\| = E(r)$;
- le plan passant par M et perpendiculaire à l'axe Ox ainsi que le plan passant par M et contenant l'axe Ox sont des plans de symétrie donc le champ appartient à ces deux plans, ce qui impose que le champ soit radial $\vec{E}(M) = E(M)\vec{u}_r$.

Finalement on a :

$$\vec{E}(M) = E(r)\vec{u}_r$$

b) Norme du champ

On peut l'obtenir très rapidement à partir du résultat obtenu pour le segment uniformément chargé en faisant tendre l vers l'infini ou α_0 vers $\frac{\pi}{2}$ et en remplaçant γ par r (pour se placer en coordonnées cylindriques). On obtient :

$$\vec{E}(M) = \frac{\lambda}{2\pi\epsilon_0 r} \vec{u}_r$$

A. Applications directes du cours

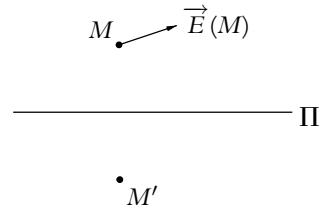
1. Champs et forces électrostatiques

Soient deux charges Q et Q' situées respectivement en M et M' , deux points distants de d .

1. Comparer les normes des champs électrostatiques auxquels sont soumises les charges Q et Q' .
2. Même question pour les forces subies par les charges.
3. Comment peut-on rendre les champs égaux en plaçant une charge Q'' sur l'axe MM' à une distance l de Q ?
4. Si Q'' est située entre les deux points M et M' , quelle est la valeur de Q'' ?

2. Plan de symétrie et plan d'antisymétrie

1. Tracer le champ électrostatique en M' , symétrique de M par rapport à Π , si Π est un plan de symétrie de la distribution de charges.
2. Même question si Π est un plan d'antisymétrie de la distribution de charges.



B. Exercices et problèmes

1. Champ au centre d'un triangle

Soient trois charges q_1 , q_2 et q_3 placées aux sommets d'un triangle équilatéral. Déterminer le champ électrostatique créé par cette distribution au centre de gravité du triangle :

1. si $q_1 = q_2 = q_3 = q$,
2. si $q_1 = q_2 = q$ et $q_3 = -q$,
3. si $q_1 = q_2 = q$ et $q_3 = 2q$,
4. si $q_1 = q$, $q_2 = 2q$ et $q_3 = -q$.

2. Electromètre

Un électromètre est constitué de deux boules métalliques identiques de masse m et de rayon r suffisamment petit pour qu'elles puissent être considérées comme ponctuelles. Elles sont suspendues à un même point O par deux fils isolants de même longueur b . Une boule notée A est fixe, le point A est sur la verticale passant par O . L'autre notée P est mobile. L'ensemble est placé dans le champ de pesanteur $\vec{g}$ supposé uniforme dans le référentiel terrestre supposé galiléen.

On donne : $b = 12 \text{ cm}$, $m = 2,55 \text{ g}$, $g = 9,81 \text{ m.s}^{-2}$ et $\frac{1}{4\pi\epsilon_0} = 9.10^9$ dans les unités du système international. Dans un premier temps, la boule P n'est pas chargée et la boule A porte

la charge électrique Q . On met les deux boules en contact. Il en résulte une déviation du fil OP d'un angle φ par rapport à la verticale.

1. Donner l'expression de l'intensité f de la force électrostatique qui s'exerce sur P .

$$1. f = \frac{1}{4\pi\epsilon_0} \frac{Q^2}{32b^2 \sin \varphi}$$

$$2. f = \frac{1}{4\pi\epsilon_0} \frac{Q^2}{8b^2 \sin 2\varphi}$$

$$3. f = \frac{1}{4\pi\epsilon_0} \frac{Q^2}{2b^2 \sin \frac{\varphi}{2}}$$

$$4. f = \frac{1}{4\pi\epsilon_0} \frac{Q^2}{16b^2 \sin^2 \frac{\varphi}{2}}$$

2. Déterminer l'expression de φ_e à l'équilibre.

$$1. \sin^2 \varphi_e = \frac{1}{4\pi\epsilon_0} \frac{Q^2}{16b^2 mg}$$

$$2. \sin^2 2\varphi_e = \frac{1}{4\pi\epsilon_0} \frac{Q^2}{8b^2 mg}$$

$$3. \sin^3 \frac{\varphi_e}{2} = \frac{1}{4\pi\epsilon_0} \frac{Q^2}{32b^2 mg}$$

$$4. \sin^4 \frac{\varphi_e}{2} = \frac{1}{4\pi\epsilon_0} \frac{Q^2}{16b^2 mg}$$

3. Déterminer l'expression de l'intensité T de la tension du fil isolant OP à l'équilibre.

$$1. T = mg$$

$$2. T = \frac{1}{4\pi\epsilon_0} \frac{Q^2}{16b^2 \sin \frac{\varphi_e}{2}} + mg \cos \varphi_e$$

$$3. T = 2mg \sin \frac{\varphi_e}{2}$$

$$4. T = \frac{1}{4\pi\epsilon_0} \frac{Q^2}{2b^2} \cos \varphi_e$$

4. En déduire la valeur de la charge Q .

$$1. 2, 0.10^{-7} \text{ C}$$

$$2. 1, 6.10^{-13} \text{ C}$$

$$3. 3, 2.10^{-16} \text{ C}$$

$$4. 4, 0.10^{-17} \text{ C}$$

5. Calculer à une constante près l'énergie potentielle de P .

$$1. Ep = \frac{1}{4\pi\epsilon_0} \frac{Q^2}{4b \sin \frac{\varphi}{2}} + mgb \cos \varphi$$

$$2. \quad E_p = \frac{1}{4\pi\epsilon_0} \frac{Q^2}{8b \sin \frac{\varphi}{2}} - mgb \cos \varphi$$

$$3. \quad E_p = \frac{1}{4\pi\epsilon_0} \frac{Q^2}{16b \cos \frac{\varphi}{2}} + mgb \sin \varphi$$

$$4. \quad E_p = \frac{1}{4\pi\epsilon_0} \frac{Q^2}{32b \sin \frac{\varphi}{2}} - mgb \cos \varphi$$

6. Qualifier l'équilibre.

1. stable
2. instable
3. indifférent

44

Propriétés du champ électrostatique

L'objet de ce chapitre est d'étudier les propriétés du champ électrique. On énonce le théorème de Gauss basé sur le flux du champ ; ce théorème fournit une méthode particulièrement efficace pour le calcul des champs électrostatiques. Ensuite, à partir de la notion de circulation d'un champ, on introduit la notion de potentiel électrostatique. On termine le chapitre par l'analyse de la topographie du champ et du potentiel électrostatique.

1. Flux du champ électrostatique Théorème de Gauss

1.1 Flux d'un champ de vecteurs

Dans ce paragraphe, on se place dans le cas général d'un champ de vecteurs quelconque qui sera noté $\vec{a}(M)$.

a) Orientation d'une surface

L'orientation d'une surface est un problème difficile : il existe des surfaces complexes pour lesquelles la définition qui va suivre ne s'applique pas. On verra alors comment procéder dans ces cas complexes.

Cas d'une surface « simple »

Soit une surface Σ s'appuyant sur une courbe fermée $\mathcal{C}$. On la suppose « simple » dans une acception qui sera précisée ultérieurement.

Pour orienter la courbe $\mathcal{C}$, on choisit un sens de parcours qui définit le sens direct.

On peut alors définir l'orientation de la surface Σ . En un point M de la courbe $\mathcal{C}$, on note $\vec{t}$ le vecteur unitaire tangent à $\mathcal{C}$ orienté dans le sens de parcours choisi.

On appelle $\vec{t}'$ le vecteur orthogonal à $\vec{t}$ dans le sens direct et situé dans le plan tangent en M à Σ . On définit alors le vecteur $\vec{n}$ normal à la surface par :

$$\vec{n} = \vec{t} \wedge \vec{t}'$$

Méthodes pratiques :

En pratique, l'orientation d'une surface se déduit de celle de la courbe par la règle suivante dite du tire-bouchon : un tire-bouchon dont on fait tourner le manche dans le sens positif de la courbe se déplace dans le sens du vecteur $\vec{n}$.

Une autre méthode consiste à placer sa main droite le long de la courbe en orientant les doigts dans le sens positif de la courbe : la direction de la surface orientée est alors donnée par la direction indiquée par le pouce.

Une dernière approche dite du « bonhomme d'Ampère » consiste à imaginer un bonhomme qui se coucherait sur la courbe en la regardant, avec la direction de la courbe qui le parcourrait des pieds à la tête : la direction de la surface orientée est donnée par celle du bras gauche tendu du bonhomme. Enfin on peut également retenir l'orientation par le schéma simple donné ci-contre.

L'orientation de la surface par le vecteur $\vec{n}$ n'est pour l'instant définie que sur la courbe $\mathcal{C}$. L'orientation de la surface en tout point s'obtient par continuité à partir de celle définie sur la courbe $\mathcal{C}$.

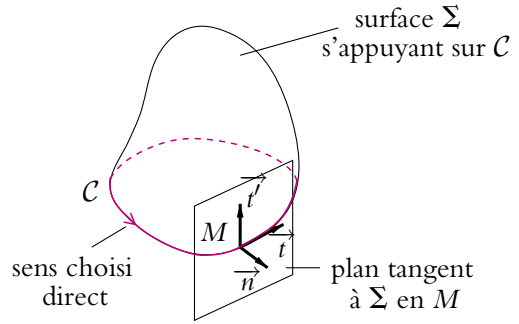


Figure 44.1 Orientation d'une surface simple.

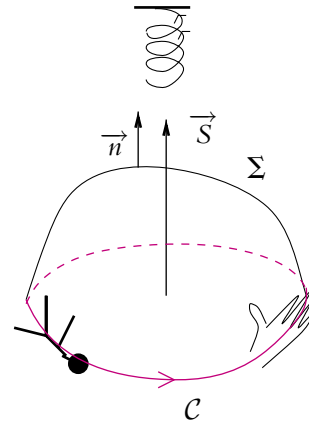


Figure 44.2 Méthodes pratiques pour orienter une surface simple.

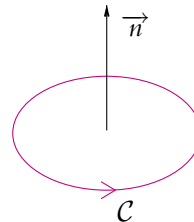


Figure 44.3 Orientation d'une surface simple.

Cas d'une surface « complexe »

Ce qui vient d'être décrit peut poser des difficultés dans le cas de surfaces « complexes ». Si on suppose que la surface n'a pas une concavité de nature constante, on peut aboutir par la méthode précédente à deux orientations contraires en un même point.

Il suffit de prendre deux courbes fermées, l'une dans la zone où la surface est concave, l'autre là où elle est convexe : par continuité, on obtient les deux orientations possibles comme le montre la figure 44.4.

Dans ce cas, on choisira une seule des deux orientations mais ce choix sera *a priori* arbitraire.

Le seul cas où l'arbitraire pourra être levé sera celui d'une surface fermée. Une surface est dite fermée si elle sépare l'espace en deux zones : le volume intérieur et le volume extérieur à la surface. On choisit habituellement l'orientation de la surface de l'intérieur vers l'extérieur.

Surface élémentaire orientée

Pour un élément de surface élémentaire dS autour de M , on appelle *surface élémentaire orientée* le vecteur :

$$d\vec{S} = dS \vec{n}$$

b) Définition d'un flux

On appelle *flux élémentaire* du champ $\vec{a}(M)$ à travers une surface élémentaire orientée $d\vec{S}_M$ la quantité :

$$d\Phi = \vec{a}(M) \cdot d\vec{S}_M$$

Le *flux* de $\vec{a}(M)$ à travers une surface Σ s'obtient en sommant les contributions des flux élémentaires associés aux surfaces élémentaires orientées :

$$\Phi = \iint_{M \in \Sigma} \vec{a}(M) \cdot d\vec{S}_M$$

c) Cas d'une surface fermée

Dans le cas d'une surface fermée, l'orientation de la surface conduit à un vecteur dirigé de l'intérieur du volume défini par la surface vers l'extérieur.

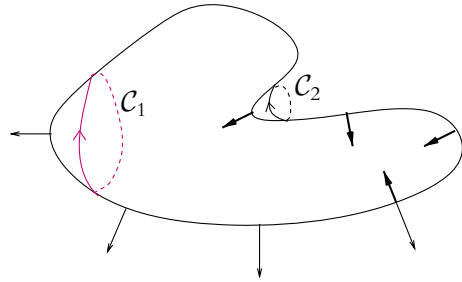


Figure 44.4 Orientation d'une surface complexe.

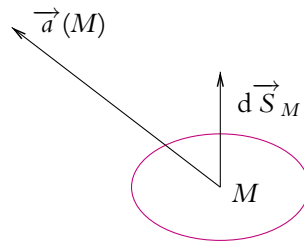


Figure 44.5 Définition du flux d'un champ de vecteurs.

De ce fait et compte-tenu des définitions précédentes, le flux à travers une surface fermée sera toujours un flux **sortant**.

1.2 Flux du champ électrostatique créé par une charge ponctuelle

a) Flux élémentaire

Soit $\vec{dS}$ une surface élémentaire autour d'un point M .

Le champ électrostatique créé en M par une charge ponctuelle q placée en O s'écrit :

$$\vec{E}(M) = \frac{1}{4\pi\epsilon_0} \frac{q}{(OM)^3} \vec{OM}$$

Le flux élémentaire de ce champ à travers $d\vec{S}$ vaut :

$$d\Phi = \vec{E}(M) \cdot d\vec{S} = \frac{1}{4\pi\epsilon_0} \frac{q}{(OM)^3} \vec{OM} \cdot d\vec{S}$$

D'après la figure ci-dessous, on a la relation :

$$\vec{OM} \cdot d\vec{S} = OM dS \cos \alpha$$

en notant α l'angle entre la droite OM et le vecteur représentant la surface élémentaire orientée $d\vec{S}$. On en déduit :

$$\frac{\vec{OM} \cdot d\vec{S}}{OM} = dS \cos \alpha = dS_0$$

où dS_0 désigne la projection de $d\vec{S}$ sur $\vec{OM}$.

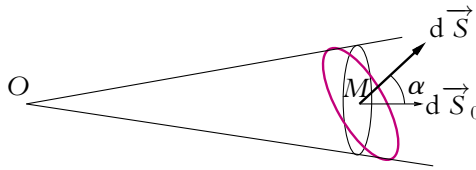


Figure 44.6 Flux élémentaire du champ électrostatique créé par une charge ponctuelle.

Le flux élémentaire du champ électrostatique créé par une charge ponctuelle q placée en O à travers n'importe quelle surface élémentaire $d\vec{S}$ au voisinage d'un point M s'exprime en fonction du flux élémentaire du champ électrostatique créé par une charge ponctuelle q placée en O à travers la surface élémentaire $d\vec{S}_0$ projection de $d\vec{S}$ sur la direction $\vec{OM}$:

$$d\Phi = \frac{1}{4\pi\epsilon_0} \frac{q}{(OM)^3} \vec{OM} \cdot d\vec{S}_0 = \frac{1}{4\pi\epsilon_0} \frac{q}{(OM)^2} dS_0$$

On notera que $\vec{dS}_0$ correspond à une surface élémentaire tracée sur la sphère de centre O passant par M .

En utilisant alors les coordonnées sphériques centrées sur O , on obtient l'expression de la surface élémentaire dS_0 vue depuis l'origine et obtenue par variation des angles θ et φ à distance constante $r = OM$ de l'origine. Cet élément de surface vaut¹ :

$$dS_0 = r^2 \sin \theta d\theta d\varphi = OM^2 \sin \theta d\theta d\varphi$$

On en déduit l'expression du flux élémentaire cherché :

$$d\Phi = \frac{1}{4\pi\epsilon_0} \frac{q}{(OM)^2} OM^2 \sin \theta d\theta d\varphi = \frac{q}{4\pi\epsilon_0} \sin \theta d\theta d\varphi$$

On remarque que cette expression est indépendante du rayon de la sphère utilisée et que seule compte l'ouverture du cône s'appuyant sur la surface élémentaire $\vec{dS}$.

b) Flux à travers une surface fermée

D'après les résultats du paragraphe précédent, le flux du champ électrostatique créé par une charge ponctuelle q placée en O à travers une surface fermée quelconque autour de la charge s'obtient en sommant les flux élémentaires à travers des surfaces élémentaires sphériques correspondant à des sphères de centre O . On a établi que le flux élémentaire est indépendant du rayon de la sphère. Par conséquent, le flux du champ électrostatique créé par une charge ponctuelle q placée en O à travers une surface fermée quelconque est le même que celui du même champ électrostatique à travers une sphère. Or la surface d'une sphère de rayon R est $4\pi R^2$ donc le flux cherché s'écrit :

$$\Phi = \frac{1}{4\pi\epsilon_0} \frac{q}{R^2} 4\pi R^2 = \frac{q}{\epsilon_0}$$

Si la surface fermée ne se trouve pas autour de la charge, le flux est nul. En effet, jusqu'à présent, on a toujours eu un produit scalaire entre $\vec{OM}$ et $\vec{dS}$ positif car $\vec{dS}_0$ est colinéaire à $\vec{OM}$ et de même sens. Ceci est lié au fait que O est à l'intérieur de la surface fermée comme le montre le schéma du bas de la figure 44.7.

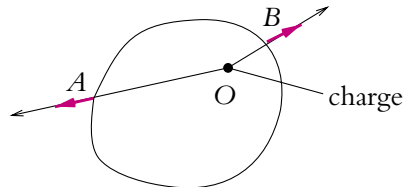
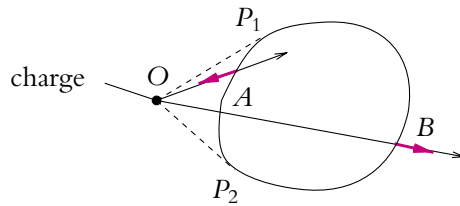


Figure 44.7 Flux du champ électrostatique à travers une surface fermée.

¹ On se reportera au paragraphe 2.3 du chapitre 42.

Si O est à l'extérieur de la surface fermée comme sur le schéma du haut de la figure 44.7 alors sur la moitié des surfaces sphériques élémentaires (partie P_1BP_2 sur le schéma), on retrouve la même situation qu'auparavant. En revanche, sur l'autre moitié (partie P_1AP_2 sur le schéma), le produit scalaire est négatif. Les deux termes correspondant à la sommation sur l'ensemble des deux situations vont donc s'annuler et le flux est nul lorsque la charge est à l'extérieur.

1.3 Théorème de Gauss

a) Exemple du fil infini uniformément chargé

On considère le cas du champ électrostatique créé par un fil infini uniformément chargé (de densité de charges linéique λ) qu'on a étudié au chapitre précédent. Son expression est en coordonnées cylindriques :

$$\vec{E}(M) = \frac{\lambda}{2\pi\epsilon_0 r} \vec{u}_r$$

Soit la surface fermée constituée d'un cylindre d'axe suivant le fil, de base circulaire de rayon r .

Le flux du champ électrostatique à travers cette surface fermée se décompose en :

$$\Phi = \Phi_{\text{lat}} + \Phi_{S_1} + \Phi_{S_2}$$

Comme $\vec{E}$ est perpendiculaire aux surfaces orientées $\vec{S}_1$ et $\vec{S}_2$, les flux Φ_{S_1} et Φ_{S_2} sont nuls. Quant au flux latéral, il s'écrit :

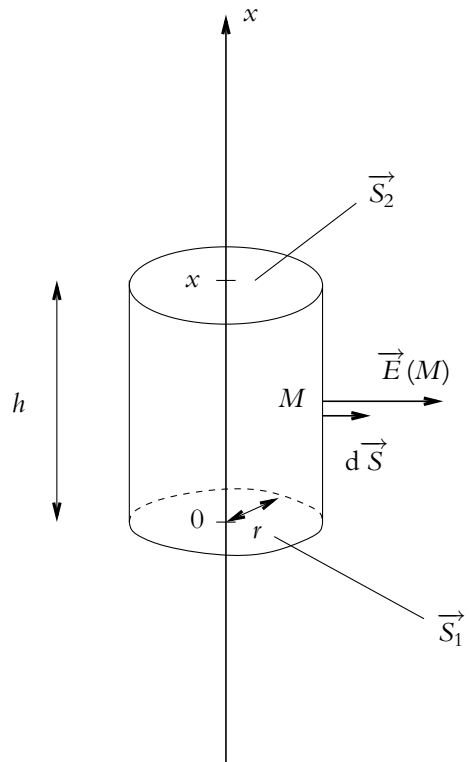
$$\Phi_{\text{lat}} = \iint_{\text{lat}} \frac{\lambda}{2\pi\epsilon_0 r} r d\theta dx$$

soit

$$\Phi_{\text{lat}} = \frac{\lambda}{2\pi\epsilon_0} \int_0^{2\pi} d\theta \int_0^h dx = \frac{\lambda h}{\epsilon_0}$$

On remarque que λh est la charge du fil contenue à l'intérieur de la surface fermée à travers laquelle on calcule le flux. On peut donc écrire le résultat précédent sous la forme :

$$\Phi = \frac{Q_{\text{int}}}{\epsilon_0}$$



On admettra la généralisation de ce résultat au cas d'une distribution de charges quelconque.

b) Enoncé pour le champ électrostatique

Le flux du champ électrostatique à travers une surface fermée $\mathcal{S}$ est égal au rapport de la charge totale contenue à l'intérieur de la surface $\mathcal{S}$ par la permittivité électrique dans le vide ϵ_0 :

$$\Phi = \oint_{M \in \mathcal{S}} \vec{E}(M) \cdot d\vec{S}_M = \frac{Q_{\text{int}}}{\epsilon_0}$$

Le « rond » autour des intégrales signifie que la surface à travers laquelle on calcule le flux est une surface fermée.

On appelle *surface de Gauss* la surface fermée $\mathcal{S}$ considérée.

Seules les charges situées dans le volume grisé sont à l'intérieur de la surface de Gauss et participent au flux du champ électrostatique à travers $\mathcal{S}$. Ce sont elles qui apparaissent dans l'énoncé du théorème de Gauss sous la notation Q_{int} .

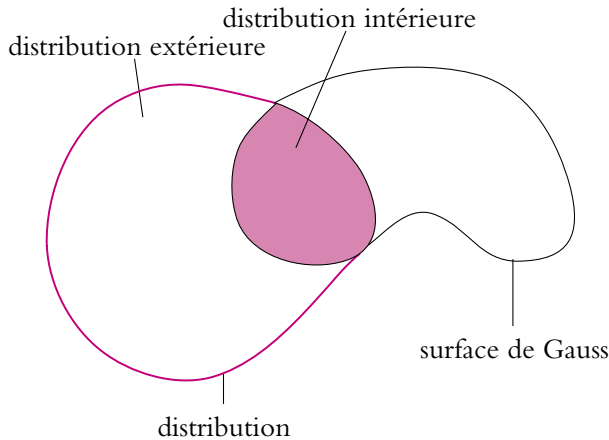


Figure 44.8 Théorème de Gauss : charges (ou masses) intérieures et extérieures.

c) Enoncé pour le champ gravitationnel

On a explicité l'analogie qui peut être établie entre le champ électrostatique et le champ gravitationnel. On peut l'utiliser ici et énoncer le théorème de Gauss pour le champ gravitationnel.

Le flux du champ gravitationnel à travers une surface fermée $\mathcal{S}$ est égal au produit de la masse totale contenue à l'intérieur de la surface $\mathcal{S}$ par -4π fois la constante de gravitation universelle G :

$$\Phi = \oint_{M \in \mathcal{S}} \vec{A}(M) \cdot d\vec{S}_M = -4\pi G M_{\text{int}}$$

En effet, par analogie, on remplace ϵ_0 par $-\frac{1}{4\pi G}$; ce qui donne le résultat annoncé.

2. Circulation du champ électrostatique - Potentiel électrostatique

2.1 Circulation d'un champ de vecteurs

Dans ce paragraphe, on se place dans le cas général d'un champ de vecteurs quelconque qui sera noté $\vec{a}(M)$.

a) Définition

Soit une courbe orientée Γ . Cela correspond à une courbe sur laquelle on a choisi un sens de parcours.

On appelle *circulation* du champ $\vec{a}(M)$ le long de Γ l'intégrale :

$$C = \int_{M \in \Gamma} \vec{a}(M) \cdot d\vec{OM}$$

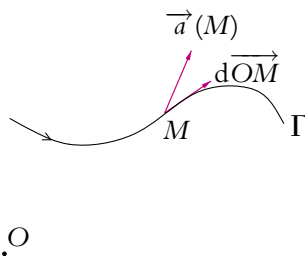


Figure 44.9 Circulation d'un vecteur.

où M parcourt la courbe Γ dans le sens de son orientation.

On note dC la circulation élémentaire, c'est-à-dire : $dC = \vec{a}(M) \cdot d\vec{OM}$.

► **Remarque** : Si le champ $\vec{a}(M)$ est un champ de forces, la circulation de ce champ n'est rien d'autre que le travail de la force le long de la courbe Γ .

b) Propriétés

- La circulation d'une somme de champs est égale à la somme des circulations des champs :

$$\int_{M \in \Gamma} \left(\sum_i \vec{a}_i(M) \right) \cdot d\vec{OM} = \sum_i \left(\int_{M \in \Gamma} \vec{a}_i(M) \cdot d\vec{OM} \right)$$

Ce résultat se démontre aisément en utilisant la linéarité de l'intégration et l'inversion de la somme finie et des intégrales.

- La circulation change de signe lorsqu'on change le sens de parcours donc l'orientation de la courbe. C'est une conséquence immédiate des propriétés de l'intégrale à savoir du changement de signe de l'intégrale lorsqu'on inverse les bornes d'intégration :

$$\int_a^b f(x)dx = - \int_b^a f(x)dx$$

- La circulation sur une courbe fermée est indépendante de l'origine choisie pour la calculer. Attention, elle n'est pas forcément nulle : elle ne le sera que si le champ dérive d'un potentiel au sens où cela sera défini au paragraphe suivant.

2.2 Notion de gradient d'une fonction scalaire

Soit f une fonction scalaire des coordonnées de l'espace et df sa différentielle.

On appelle *gradient de la fonction* f et on note $\overrightarrow{\text{grad}}f$ la quantité vérifiant :

$$\overrightarrow{\text{grad}}f \cdot d\overrightarrow{OM} = df$$

On cherche maintenant à exprimer le gradient de f dans les différents systèmes de coordonnées, en explicitant le produit scalaire en fonction des déplacements élémentaires exprimés dans chaque système de coordonnées.

a) Coordonnées cartésiennes

On a :

$$\begin{aligned} df &= \left(\overrightarrow{\text{grad}}f \right)_x dx + \left(\overrightarrow{\text{grad}}f \right)_y dy + \left(\overrightarrow{\text{grad}}f \right)_z dz \\ &= \left(\frac{\partial f}{\partial x} \right)_{y,z} dx + \left(\frac{\partial f}{\partial y} \right)_{x,z} dy + \left(\frac{\partial f}{\partial z} \right)_{x,y} dz \end{aligned}$$

On note en indice les variables supposées constantes au cours de la dérivation : par exemple, $\left(\frac{\partial f}{\partial x} \right)_{y,z}$ est la dérivée partielle de f en tant que fonction de x uniquement, les variables y et z étant supposés fixes. Par identification, on obtient :

$$\overrightarrow{\text{grad}}f = \left(\frac{\partial f}{\partial x} \right)_{y,z} \overrightarrow{u_x} + \left(\frac{\partial f}{\partial y} \right)_{x,z} \overrightarrow{u_y} + \left(\frac{\partial f}{\partial z} \right)_{x,y} \overrightarrow{u_z}$$

b) Coordonnées cylindriques

On a :

$$\begin{aligned} df &= \left(\overrightarrow{\text{grad}}f \right)_r dr + \left(\overrightarrow{\text{grad}}f \right)_\theta r d\theta + \left(\overrightarrow{\text{grad}}f \right)_z dz \\ &= \left(\frac{\partial f}{\partial r} \right)_{\theta,z} dr + \left(\frac{\partial f}{\partial \theta} \right)_{r,z} d\theta + \left(\frac{\partial f}{\partial z} \right)_{r,\theta} dz \end{aligned}$$

soit par identification :

$$\overrightarrow{\text{grad}f} = \left(\frac{\partial f}{\partial r}\right)_{\theta,z} \overrightarrow{u_r} + \frac{1}{r} \left(\frac{\partial f}{\partial \theta}\right)_{r,z} \overrightarrow{u_\theta} + \left(\frac{\partial f}{\partial z}\right)_{r,\theta} \overrightarrow{u_z}$$

c) Coordonnées sphériques

On a

$$\begin{aligned} df &= \left(\overrightarrow{\text{grad}f}\right)_r dr + \left(\overrightarrow{\text{grad}f}\right)_\theta r d\theta + \left(\overrightarrow{\text{grad}f}\right)_\varphi r \sin \theta d\varphi \\ &= \left(\frac{\partial f}{\partial r}\right)_{\theta,\varphi} dr + \left(\frac{\partial f}{\partial \theta}\right)_{r,\varphi} d\theta + \left(\frac{\partial f}{\partial \varphi}\right)_{r,\theta} d\varphi \end{aligned}$$

soit par identification :

$$\overrightarrow{\text{grad}f} = \left(\frac{\partial f}{\partial r}\right)_{\theta,\varphi} \overrightarrow{u_r} + \frac{1}{r} \left(\frac{\partial f}{\partial \theta}\right)_{r,\varphi} \overrightarrow{u_\theta} + \frac{1}{r \sin \theta} \left(\frac{\partial f}{\partial \varphi}\right)_{r,\theta} \overrightarrow{u_\varphi}$$

d) Interprétation physique

La notion de gradient permet de quantifier la variation spatiale d'une grandeur. Par exemple, une variation de température ou de densité selon un axe se traduit par un gradient.

2.3 Potentiel électrostatique créé par une charge ponctuelle

a) Circulation du champ électrostatique créé par une charge ponctuelle

Soit une charge ponctuelle q placée en un point P .

Le champ électrostatique créé par cette charge en un point M s'écrit :

$$\overrightarrow{E}(M) = \frac{1}{4\pi\epsilon_0} \frac{q}{(PM)^3} \overrightarrow{PM}$$

La circulation élémentaire de ce champ s'exprime par :

$$dC = \overrightarrow{E}(M) \cdot d\overrightarrow{OM} = \frac{1}{4\pi\epsilon_0} \frac{q}{(PM)^3} \overrightarrow{PM} \cdot d\overrightarrow{OM}$$

P étant un point fixe, $d\overrightarrow{OP} = \overrightarrow{0}$ et

$$d\overrightarrow{OM} = d\overrightarrow{OP} + d\overrightarrow{PM} = d\overrightarrow{PM}$$

Or $\overrightarrow{PM} \cdot d\overrightarrow{PM} = \frac{1}{2} d(\overrightarrow{PM} \cdot \overrightarrow{PM}) = \frac{1}{2} d(PM^2) = PM dPM$. On en déduit :

$$dC = \frac{q}{4\pi\epsilon_0} \frac{d(PM)}{(PM)^2}$$

Or

$$\frac{d(PM)}{(PM)^2} = -d\left(\frac{1}{PM}\right)$$

d'où :

$$dC = -d\left(\frac{1}{4\pi\epsilon_0} \frac{q}{PM}\right)$$

On en déduit : $dC = -dV$ avec

$$V = \frac{1}{4\pi\epsilon_0} \frac{q}{PM} + \text{constante}$$

et, le long d'une courbe Γ , $C = - \int_{M \in \Gamma} d\left(\frac{1}{4\pi\epsilon_0} \frac{q}{PM}\right)$.

b) Potentiel électrostatique créé par une charge ponctuelle

Au paragraphe précédent, on a établi que la circulation élémentaire du champ électrostatique s'écrit comme l'opposé de la différentielle d'une quantité scalaire que l'on a notée V . La grandeur V est appelée *potentiel électrostatique*.

On a la relation :

$$dC = \vec{E}(M) \cdot d\vec{OM} = -dV(M) \quad \text{avec} \quad V(M) = \frac{1}{4\pi\epsilon_0} \frac{q}{PM} + \text{constante}$$

Le potentiel électrostatique $V(M)$ n'est défini, d'après ce qui précède, qu'à une constante près. Le plus souvent, on la choisit nulle à l'infini c'est-à-dire quand on est infiniment loin de la charge créant le champ électrostatique. Cela revient à dire que deux charges infiniment éloignées l'une de l'autre n'exercent aucune action électrostatique sur l'autre. Dans le cas où on a des distributions d'extension finie, ce choix est donc parfaitement justifié. En revanche, si on considère une distribution d'extension infinie, cela signifie que l'on doit tenir compte de l'interaction de charges infiniment éloignées. Il faudra donc recourir à un autre choix de constante pour le potentiel. Ici la distribution est finie et on obtient :

$$V(M) = \frac{1}{4\pi\epsilon_0} \frac{q}{PM}$$

Un champ vérifiant la propriété $dC = -dV$ est dit *champ de gradient* et on peut écrire :

$$\vec{E} = -\vec{\text{grad}} V$$

En effet, on a établi : $\vec{E}(M) \cdot d\vec{OM} = -dV$ et par définition du gradient : $dV = \vec{\text{grad}} V \cdot d\vec{OM}$.

2.4 Potentiel électrostatique d'une distribution quelconque

Ce qui vient d'être établi pour une charge ponctuelle se généralise au cas d'une distribution de charges quelconque.

a) Exemple du fil infini uniformément chargé

On reprend l'exemple du fil infini pour lequel on a établi l'expression du champ électrostatique qu'il crée en un point M au chapitre précédent. On rappelle son expression en coordonnées cylindriques :

$$\vec{E}(M) = \frac{\lambda}{2\pi\epsilon_0 r} \vec{u}_r$$

Le calcul de la circulation C du champ électrostatique le long de Γ s'écrit en coordonnées cylindriques :

$$C = \int_{M \in \Gamma} \vec{E}(M) \cdot d\vec{OM} = \int_{M \in \Gamma} \frac{\lambda}{2\pi\epsilon_0 r} dr = \frac{\lambda}{2\pi\epsilon_0} \ln\left(\frac{r}{r_0}\right)$$

$$\text{donc } V(M) = \frac{\lambda}{2\pi\epsilon_0} \ln\left(\frac{r}{r_0}\right)$$

En choisissant l'origine des potentiels pour $r = r_0$ et non à l'infini : la distribution est infinie et il y a donc des charges à l'infini, ce qui interdit de prendre l'origine des potentiels à l'infini.

b) Distributions ponctuelles

Dans ce cas, il suffit d'appliquer le principe de superposition :

$$C = \int_{M \in \mathcal{C}} \vec{E}(M) \cdot d\vec{OM} = \int_{M \in \mathcal{C}} \left(\sum_i \vec{E}_i(M) \right) \cdot d\vec{OM} = \sum_i \int_{M \in \mathcal{C}} \vec{E}_i(M) \cdot d\vec{OM}$$

Et pour chaque champ $\vec{E}_i(M)$, on utilise le même procédé que pour une charge ponctuelle unique :

$$\int_{M \in \mathcal{C}} \vec{E}_i(M) \cdot d\vec{OM} = - \int_{M \in \mathcal{C}} \frac{q_i}{4\pi\epsilon_0} d\left(\frac{1}{P_i M}\right)$$

soit finalement :

$$C = - \sum_i \int_{M \in \mathcal{C}} \frac{q_i}{4\pi\epsilon_0} d\left(\frac{1}{P_i M}\right) = - \int_{M \in \mathcal{C}} d\left(\frac{1}{4\pi\epsilon_0} \sum_i \frac{q_i}{P_i M}\right)$$

d'où, pour la circulation élémentaire : $dC = -dV$ avec le potentiel électrostatique

$$V(M) = \frac{1}{4\pi\epsilon_0} \sum_i \frac{q_i}{P_i M} + \text{constante}$$

et

$$\vec{E} = -\vec{\text{grad}} V$$

c) Distributions continues

Tout ce qui est exposé dans ce paragraphe ne concerne que les distributions d'extension finie pour lesquelles on choisit l'origine des potentiels à l'infini. On a vu sur l'exemple du fil infini uniformément chargé qu'il est nécessaire d'utiliser la définition du potentiel à partir de la circulation du champ électrostatique pour les distributions d'extension infinie.

Pour les distributions continues et d'**extension finie**, on obtient l'expression du potentiel électrostatique avec origine des potentiels à l'infini par une démonstration analogue à celle faite pour les distributions discrètes. C'est l'absence de charges à l'infini qui rend possible le choix de l'origine des potentiels à l'infini. On se contente donc de donner les expressions relatives à chacune des distributions continues de charges :

- distribution volumique :

$$V(M) = \frac{1}{4\pi\epsilon_0} \iiint_{P \in \mathcal{V}} \frac{\rho(P)}{PM} d\tau_P$$

- distribution surfacique :

$$V(M) = \frac{1}{4\pi\epsilon_0} \iint_{P \in \mathcal{S}} \frac{\sigma(P)}{PM} dS_P$$

- distribution linéique :

$$V(M) = \frac{1}{4\pi\epsilon_0} \int_{P \in \mathcal{L}} \frac{\lambda(P)}{PM} dl_P$$

Dans tous les cas,

$$\vec{E} = -\vec{\text{grad}} V$$

On note qu'il s'agit ici d'intégrales scalaires contrairement à ce qu'on avait pour le calcul du champ où l'intégration est vectorielle et où il fallait projeter pour obtenir des intégrales scalaires.

Lorsque la distribution est infinie, il convient de changer l'origine des potentiels. L'expression obtenue dans le cas d'une charge ponctuelle (et utilisée ici pour les distributions d'extension finie) n'est plus valable. Les intégrales précédentes qui découlent de l'application du principe de superposition n'opèrent pas sur des expressions valables et n'ont donc plus de raison d'être. On notera que, mathématiquement, ce problème se traduit par le fait qu'il n'y a *a priori* aucune garantie que les intégrales précédentes soient convergentes lorsque le point P tend vers l'infini.

d) Circulation du champ électrostatique le long d'un contour fermé

Du fait que le champ électrostatique est un champ de gradient, on en déduit immédiatement que sa circulation le long d'un contour fermé est nulle. En effet, on a :

$$\oint_{M \in \mathcal{C}} \vec{E}(M) \cdot d\vec{OM} = - \oint_{M \in \mathcal{C}} dV(M) = V(M_i) - V(M_f) = 0$$

car $M_i = M_f$ du fait du caractère fermé de $\mathcal{C}$.

Le « rond » autour de l'intégrale signifie que la circulation est calculée le long d'une courbe fermée.

2.5 Continuité du potentiel

Le problème de continuité peut se poser lorsque l'intégrale qui permet son calcul n'est pas définie. Ceci risque de se produire dans deux cas :

- lorsque la distribution est d'extension infinie, ce dont on ne s'est pas préoccupé ici,
- lorsque PM tend vers zéro : le terme $\frac{1}{PM}$ tend vers l'infini et on n'est pas sûr *a priori* que l'intégrale soit définie.

Dans la suite de ce paragraphe, on exclut les distributions d'extension infinie : on n'a donc pas de charges à l'infini. Cela se traduit mathématiquement par le fait que :

$$\lim_{M \rightarrow \infty} \rho(M) = 0$$

Le seul problème qui peut se poser est celui d'un point M appartenant à la distribution soit

$$\lim_{M \rightarrow P} \frac{1}{PM} = \infty$$

Dans le cas d'un volume, l'élément de volume exprimé en coordonnées sphériques vaut : $r^2 \sin \theta dr d\theta d\varphi$ d'où l'intégrale s'écrit :

$$\begin{aligned} V(M) &= \frac{1}{4\pi\epsilon_0} \iiint_{P \in \mathcal{V}} \frac{\rho(P)}{r} r^2 \sin \theta dr d\theta d\varphi \\ &= \frac{1}{4\pi\epsilon_0} \iiint_{P \in \mathcal{V}} \rho(P) r \sin \theta dr d\theta d\varphi \end{aligned}$$

et l'intégrale est convergente.

De même, dans le cas d'une surface, l'élément de surface exprimé en coordonnées cylindriques vaut : $r dr d\theta$ et l'intégrale s'écrit :

$$V(M) = \frac{1}{4\pi\epsilon_0} \iint_{P \in \mathcal{S}} \frac{\sigma(P)}{r} r dr d\theta = \frac{1}{4\pi\epsilon_0} \iint_{P \in \mathcal{S}} \sigma(P) dr d\theta$$

qui est également une intégrale convergente.

En revanche, pour un fil, l'élément de longueur s'écrit dr et l'intégrale :

$$V(M) = \frac{1}{4\pi\epsilon_0} \int_{P \in \mathcal{L}} \frac{\lambda(P)}{r} dr$$

est divergente pour r tendant vers 0. Cela se traduit par une discontinuité du potentiel au niveau du fil.

Le potentiel est donc continu pour un volume chargé ou une surface chargée mais présente des discontinuités pour un fil chargé. Cela correspond physiquement au fait que le champ électrostatique n'est pas défini sur le fil.

2.6 Conséquences des symétries des distributions sur le potentiel électrostatique

Soit $\mathcal{D}$ une distribution volumique de charges de densité ρ créant un potentiel électrostatique $V(M)$ au point M . Par la symétrie par rapport à un plan $\mathcal{P}$, le point M devient le point M' , la distribution $\mathcal{D}$ la distribution $\mathcal{D}'$ de densité ρ' et le potentiel $V'(M')$, créé par $\mathcal{D}'$, s'écrit du fait de l'invariance des lois de l'électrostatique :

$$V'(M') = \frac{1}{4\pi\epsilon_0} \iiint_{P' \in \mathcal{D}'} \frac{\rho'(P') d\tau_{P'}}{P'M'}$$

La transformation étant une isométrie, on a $P'M' = PM$.

Si la distribution $\mathcal{D}$ est symétrique par rapport à $\mathcal{P}$, $\rho'(P') = \rho(P)$ et P (ou P') décrit $\mathcal{D}$ ou $\mathcal{D}'$, c'est-à-dire qu'on peut considérer indifféremment que le potentiel est créé par $\mathcal{D}$ ou par $\mathcal{D}'$. On en déduit

$$V(M) = V(M')$$

Si la distribution $\mathcal{D}$ est antisymétrique par rapport à $\mathcal{P}$, $\rho'(P') = -\rho(P)$ et P (ou P') décrit $\mathcal{D}$ ou $\mathcal{D}'$, c'est-à-dire qu'on peut considérer indifféremment que le potentiel est créé par $\mathcal{D}$ ou $\mathcal{D}'$. On en déduit

$$V(M) = -V(M')$$

3. Énergie électrostatique

3.1 Énergie potentielle d'interaction

a) Énergie potentielle d'une charge ponctuelle dans un champ électrostatique extérieur fixe

Soit une charge q placée en un point M .

On suppose que dans l'espace règne un champ électrostatique $\vec{E}(M)$ dérivant du potentiel $V(M)$.

Lors d'un déplacement $d\overrightarrow{OM} = \overrightarrow{MM'}$ du point M au point M' de la charge q , la charge q est soumise à la force électrostatique $\overrightarrow{F}$ du fait de l'existence du champ électrostatique $\overrightarrow{E}(M)$. Elle subit le travail :

$$\delta W = \overrightarrow{F} \cdot d\overrightarrow{OM} = q \overrightarrow{E}(M) \cdot d\overrightarrow{OM}$$

Or $\overrightarrow{E}(M) = -\overrightarrow{\text{grad}}V(M)$ et par définition du gradient :

$$\overrightarrow{E}(M) \cdot d\overrightarrow{OM} = -dV(M)$$

D'où :

$$\delta W = -q dV(M) = -d(qV(M))$$

Le travail de la force électrostatique entre deux positions M_1 et M_2 s'écrit :

$$W = q(V(M_1) - V(M_2)) = q(V_1 - V_2)$$

où V_1 et V_2 sont respectivement les potentiels aux positions initiale et finale. On constate que ce travail ne dépend pas du chemin suivi mais uniquement des positions finale et initiale : la force est une force conservative.

Elle dérive d'une énergie potentielle E_p telle que :

$$\delta W = -dE_p$$

En utilisant l'expression trouvée précédemment, on en déduit :

$$E_p = qV + \text{constante}$$

E_p est l'énergie potentielle de la charge q dans le potentiel électrostatique $V(M)$ dont dérive le champ électrostatique $\overrightarrow{E}(M)$.

► Remarque : gradient et forces conservatives

Le gradient fournit un outil mathématique permettant de donner une nouvelle définition des forces conservatives. On dit qu'une force dérive d'un potentiel ou qu'elle est conservative s'il existe une fonction U telle que

$$\overrightarrow{f} = -\overrightarrow{\text{grad}}U$$

La quantité U est appelée énergie potentielle.

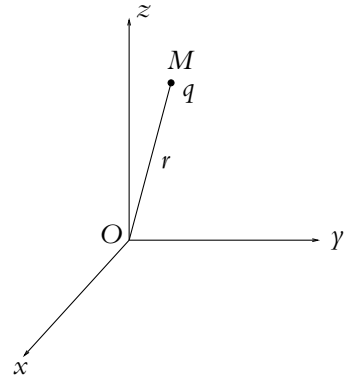


Figure 44.10 Charge q dans un champ $\overrightarrow{E}$.

Le travail d'une force ainsi définit s'obtient par :

$$\delta W = \vec{f} \cdot d\vec{OM} = -\vec{\text{grad}}U \cdot d\vec{OM} = -dU$$

par définition du gradient.

On retrouve que le travail d'une force conservative au cours d'un déplacement est égal à l'opposé de la variation d'énergie potentielle. De même, par le théorème de l'énergie cinétique, on exprime la conservation de l'énergie mécanique pour une force conservative. Il s'agit de la même notion que celle déjà traitée en mécanique sans avoir fait usage de la notion de gradient.

b) Énergie potentielle d'interaction entre deux charges

On applique ce résultat au cas où le champ est créé par une charge ponctuelle q' placée en P . Le potentiel créé en un point M par cette charge q' située en P s'exprime par :

$$V(M) = \frac{1}{4\pi\epsilon_0} \frac{q'}{PM}$$

L'énergie potentielle de la charge q placée en M vaut alors :

$$E_p = qV(M) = \frac{1}{4\pi\epsilon_0} \frac{qq'}{PM}$$

Il s'agit de l'énergie potentielle d'interaction entre les deux charges q et q' placées respectivement en M et P .

► **Remarque** : Cette énergie potentielle est également l'énergie potentielle de la charge q' placée en P dans le champ électrostatique créé par la charge q située en M :

$$E_{p'} = q'V'(P) = \frac{1}{4\pi\epsilon_0} \frac{qq'}{PM} = E_p$$

On obtient une expression parfaitement symétrique vis-à-vis de q et q' , ce qui justifie parfaitement la dénomination employée d'*énergie potentielle d'interaction*.

3.2 Énergie électrostatique

a) Définition

On appelle *énergie électrostatique* d'une distribution de charges l'énergie nécessaire pour construire cette distribution de manière réversible en « prenant les charges à l'infini ».

Il convient de bien noter qu'il s'agit de l'énergie juste nécessaire à la constitution de la distribution d'un point de vue électrostatique. On exclut toute autre forme d'interaction.

D'autre part, quand on dit que les charges sont initialement à l'infini, cela signifie que toutes les charges de la distribution sont initialement infiniment éloignées les unes des autres, de sorte qu'il n'y a pas d'interaction électrostatique entre elles. Cela correspond à prendre une énergie potentielle nulle à l'infini.

On peut effectuer cette opération en apportant de l'énergie soit de façon mécanique par un opérateur extérieur, soit de façon électrique par apport de charges à l'aide de générateurs.

Le calcul de cette énergie nécessite d'évaluer le travail fourni par l'opérateur extérieur ou par les générateurs pour placer les charges dans leur état actuel.

b) Bilan énergétique et interprétation de cette définition

Soit une particule de masse m de charge q dont on étudie le mouvement.

Système : la charge q de masse m .

Référentiel : celui du laboratoire considéré comme galiléen.

Bilan des forces :

- force d'origine électrostatique,
- poids qui sera négligé ici car il est toujours négligeable devant les forces électrostatiques (on précisera ce point dans le chapitre sur le mouvement des particules chargées),
- action de l'opérateur qui déplace la charge q .

Le bilan énergétique peut être obtenu à partir du théorème de l'énergie cinétique : la variation d'énergie cinétique est égale à la somme des travaux des forces soit

$$dE_c = \delta W_{\text{elec}} + \delta W_{\text{op}}$$

Or on vient d'établir que $\delta W_{\text{elec}} = -dEp$, on en déduit :

$$dEp = -dE_c + \delta W_{\text{op}}$$

En intégrant entre l'instant où la charge n'est soumise à aucune action électrostatique et l'instant où elle est dans sa position finale (ou actuelle), on a, puisque $\Delta Ep = Ep - 0$:

$$Ep = -\Delta E_c + W_{\text{op}}$$

D'autre part, il s'agit de construire la distribution et non de communiquer aux charges une énergie cinétique. Initialement les charges sont immobiles et après leur transfert depuis l'infini à leur position dans la distribution, elles le sont à nouveau. On a donc :

$$\Delta E_c = 0$$

On interprète donc la définition de l'énergie électrostatique par la relation :

$$Ep = W_{\text{op}}$$

Il s'agit donc du travail que doit fournir l'opérateur pour constituer la distribution à partir de charges infiniment éloignées.

► Remarques

1. L'opérateur doit exercer une force opposée à la force électrostatique pour amener la charge dans sa position actuelle :

$$\delta W_{\text{op}} = dEp = -\delta W_{\text{élec}}$$

par définition de l'énergie potentielle électrostatique et calcul du travail de la force électrostatique. Donc

$$\vec{F}_{\text{op}} \cdot d\vec{OM} = -\vec{F}_{\text{élec}} \cdot d\vec{OM}$$

pour tout déplacement $d\vec{OM}$. On en déduit :

$$\vec{F}_{\text{op}} = -\vec{F}_{\text{élec}}$$

On peut également obtenir ce résultat par application du principe des actions réciproques en considérant les charges et l'opérateur ponctuels.

2. L'opérateur exerce une force pour maintenir la charge à sa position actuelle : si ce n'était pas le cas, la charge ne serait soumise qu'à la force électrostatique et le principe fondamental impliquerait un mouvement de la charge puisque la somme des forces et donc son accélération ne seraient pas nulles.
3. La construction de la distribution est quasi-statique et infiniment lente au sens où cela est défini en thermodynamique.

3.3 Énergie électrostatique propre d'un système de deux charges

On s'intéresse à l'énergie électrostatique propre d'un système de charges ponctuelles c'est-à-dire à l'énergie potentielle d'interaction entre les charges constituant le système, en se limitant au cas de deux charges. La généralisation au cas de N charges sera étudiée en exercice.

Soient deux charges ponctuelles q_1 et q_2 placées en M_1 et M_2 .

L'énergie électrostatique est égale au travail fourni par un opérateur pour construire cette distribution en déplaçant les charges q_1 et q_2 depuis une position initiale où elles n'exercent aucune interaction électrostatique l'une sur l'autre : on les suppose classiquement situées à l'infini, c'est-à-dire infiniment éloignées de sorte que l'interaction électrostatique peut être négligée.

On commence par amener la charge q_1 depuis l'infini jusqu'à sa position finale. Il ne règne aucun champ électrostatique : l'opérateur n'a pas à opposer de force à une force électrostatique inexistante. Il n'y a donc pas de travail de l'opérateur et la contribution à l'énergie électrostatique est nulle.

On amène ensuite la charge q_2 depuis l'infini. Elle va être soumise au champ électrostatique créé par q_1 déjà positionnée. L'opérateur doit exercer une force pour s'opposer à la force due à la charge q_1 et donc un travail :

$$\delta W_{\text{op}} = dEp = -\delta W$$

avec :

$$dEp = q_2 dV_2 \quad \text{et} \quad V_2 = \frac{1}{4\pi\epsilon_0} \frac{q_1}{M_1 M_2}$$

D'où :

$$Ep = \frac{1}{4\pi\epsilon_0} \frac{q_1 q_2}{M_1 M_2}$$

A priori, il faudrait ajouter une constante mais celle-ci sera prise nulle à l'infini et disparaît dans l'expression de l'énergie potentielle.

On écrit cette énergie potentielle sous la forme :

$$Ep = q_2 V_2$$

en notant V_2 le potentiel créé par les autres charges au point où se trouve q_2 (en l'occurrence créé par q_1).

On peut également construire cette distribution en amenant la charge q_2 en premier. Cette opération ne nécessite aucune énergie : il n'y a ni champ ni potentiel électrostatiques. Le fait d'amener ensuite q_1 nécessite de la part de l'opérateur de fournir un travail :

$$W_{\text{op}} = q_1 V_1$$

en notant V_1 le potentiel créé par les autres charges au point où se trouve q_1 (en l'occurrence créé par q_2). On procède de la même manière que pour la construction précédente.

Donc

$$Ep = q_1 V_1 = q_2 V_2$$

On peut donc symétriser l'expression en notant l'énergie électrostatique

$$Ep = \frac{1}{2} (q_1 V_1 + q_2 V_2)$$

Il s'agit de l'énergie électrostatique propre du système formé par les deux charges q_1 et q_2 placées respectivement en M_1 et en M_2 .

4. Topographies du champ et du potentiel électrostatiques

4.1 Topographie du champ électrostatique

Lors de l'étude d'un champ, plusieurs types de courbes fournissent une représentation des caractéristiques du champ. On les définit ici dans le cas du champ électrostatique mais cela restera valable dans tous les cas où on envisagera un champ vectoriel comme le champ gravitationnel, le champ magnétostatique ou un champ de vitesse en mécanique des fluides (Cf. cours de seconde année option PC/PC* et PSI/PSI*).

4.2 Lignes de champ

On appelle *ligne de champ* une courbe tangente en chaque point au vecteur champ électrostatique.

Soit M un point où on cherche la ligne de champ de $\vec{E}(M)$. On note O l'origine du référentiel dans lequel on se place. La direction de la tangente à la courbe sera donnée par le vecteur $d\vec{OM}$, vecteur déplacement élémentaire du point M .

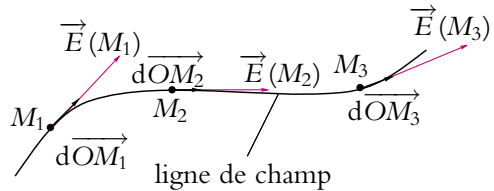


Figure 44.11 Définition d'une ligne de champ.

La définition de la ligne de champ impose que ce vecteur $d\vec{OM}$ soit colinéaire à $\vec{E}(M)$, ce qui peut se traduire par :

$$d\vec{OM} = k(M) \vec{E}(M)$$

où $k(M)$ est un réel dépendant du point M , ou par :

$$\vec{E}(M) \wedge d\vec{OM} = \vec{0}$$

La ligne de champ se réduit à un point dans le cas où le champ est nul.

4.3 Equipotentielles

On appelle *équipotentielle* la surface reliant l'ensemble des points ayant la même valeur du potentiel électrostatique à savoir

$$V(M) = \text{constante} \quad \text{ou} \quad dV(M) = V(M) - V(M') = 0$$

pour M et M' deux points très proches sur l'équipotentielle.

Lors des représentations qui seront faites, on se placera souvent dans un plan. Les équipotentielles se traduiront par des courbes qui sont les traces de la surface dans un plan.

4.4 Propriétés des lignes de champ et des équipotentielles

a) Le potentiel décroît le long d'une ligne de champ

À partir de la relation : $dV(M) = -\vec{E}(M) \cdot d\vec{OM}$, on en déduit que, le long d'une ligne de champ, la variation du potentiel est négative. En effet, $d\vec{OM}$ est par définition de la ligne de champ orienté dans le même sens que le champ $\vec{E}(M)$ donc le produit scalaire de l'un par l'autre est positif.

b) Convergence et divergence des lignes de champ

Si le potentiel admet un extremum en un point, alors le champ électrostatique en ce point est nul et les lignes de champ en ce point peuvent prendre des directions quelconques (le vecteur nul est colinéaire à toute courbe).

En utilisant le résultat du paragraphe précédent, on obtient :

- si le potentiel est maximal en un point, les lignes de champ divergent de ce point,
- si le potentiel est minimal en un point, les lignes de champ convergent vers ce point.

c) Le potentiel n'admet pas d'extrema en dehors des charges

Le flux du champ électrostatique à travers une surface fermée entourant un point où le potentiel électrostatique admet un extremum est non nul du fait de la convergence ou de la divergence des lignes de champ.

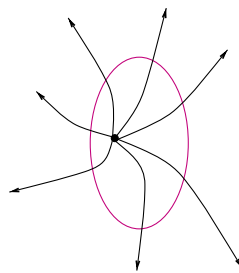
En appliquant le théorème de Gauss :

$$\Phi = \frac{Q_{\text{int}}}{\epsilon_0}$$

on en déduit que la charge intérieure est non nulle.

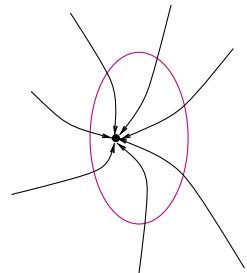
- Si $Q_{\text{int}} > 0$, les lignes de champ sont orientées vers l'extérieur de la surface de Gauss et le champ étant orienté dans le sens des potentiels décroissants, le potentiel est maximal.

- Si $Q_{\text{int}} < 0$, les lignes de champ sont orientées vers l'intérieur de la surface de Gauss et le champ étant orienté dans le sens des potentiels décroissants, le potentiel est minimal.



$$Q_{\text{int}} > 0$$

V maximal



$$Q_{\text{int}} < 0$$

V minimal

Figure 44.12 Extrema du potentiel électrostatique.

d) Orthogonalité des lignes de champ et des équipotentielles

D'après la définition du gradient, on a :

$$dV(M) = \overrightarrow{\text{grad}} V(M) \cdot d\overrightarrow{OM} = -\overrightarrow{E}(M) \cdot d\overrightarrow{OM}$$

Si le point M appartient à une équipotentielle alors le vecteur $d\overrightarrow{OM}$ est tangent à la surface de l'équipotentielle et $dV(M) = 0$. La relation :

$$dV(M) = -\overrightarrow{E}(M) \cdot d\overrightarrow{OM} = 0$$

implique que les équipotentielles sont les surfaces perpendiculaires en tout point au champ électrostatique qui règne en ce point. On en déduit que dans un plan contenant les lignes de champ, ces dernières (colinéaires au champ) et les courbes équipotentielles, traces des surfaces équipotentielles dans le plan des lignes de champ (perpendiculaires au champ), sont des courbes orthogonales.

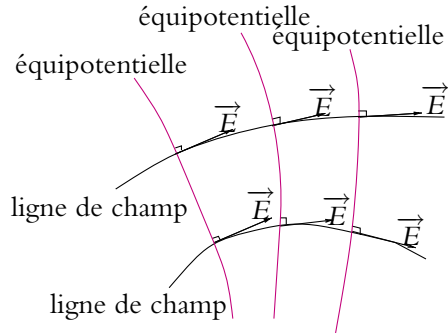


Figure 44.13 Equipotentielles et lignes de champ.

4.5 Exemple de topographie d'un champ

Soit la distribution discrète constituée de quatre charges $q_1 = e$, $q_2 = 3e$, $q_3 = -3e$ et $q_4 = -e$ situées respectivement aux points de coordonnées cartésiennes dans le plan $A_1(-1, 0)$, $A_2(1, 0)$, $A_3(0, -1)$ et $A_4(0, 1)$. Un logiciel de simulation permet d'obtenir la carte des lignes de champ et des équipotentielles de cette distribution. Elles sont données par les figures 44.14 et 44.15.

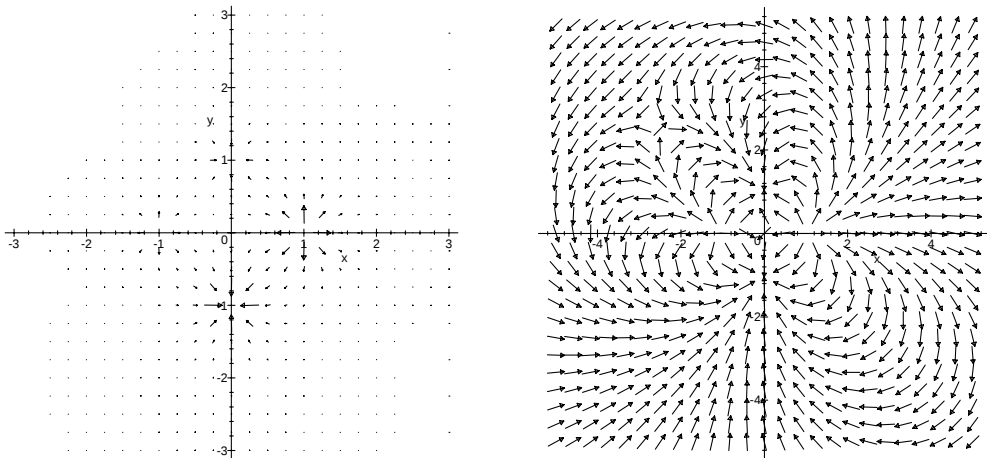


Figure 44.14 Lignes de champ d'une distribution discrète de quatre charges. La longueur des flèches est proportionnelle à la norme du champ sur la figure de gauche ; sur la figure de droite, les flèches ont même longueur indépendamment de la norme du champ.

On constate que les lignes de champ divergent d'un point où se trouve une charge négative et qu'elles convergent vers un point où se trouve une charge positive conformément aux propriétés précédemment énoncées. On peut ainsi à partir d'une carte de lignes de champ déterminer les points où il y a une charge et connaître le signe de cette dernière.

On remarque également sur la figure de gauche où la longueur des flèches est proportionnelle à la norme du champ que la norme du champ décroît avec la distance aux charges qui le créent. En effet, le champ est plus important à proximité des charges.

Enfin, le plan d'équation $x = -y$ est un plan d'antisymétrie de la distribution de charges, le champ électrique est bien perpendiculaire à ce plan en un point de ce plan.

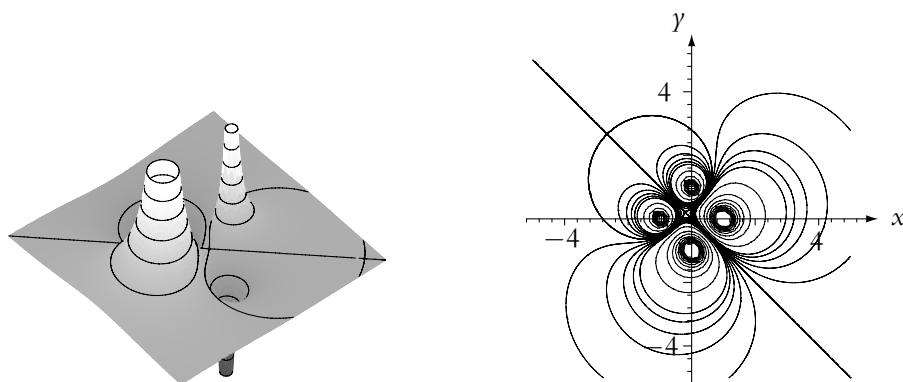


Figure 44.15 Équipotentielles d'une distribution discrète de quatre charges. Le niveau de gris traduit l'intensité relative du potentiel.

On note que les équipotentielles sont plus resserrées au voisinage des charges c'est-à-dire là où le champ est plus intense (Cf. remarques précédentes sur l'analyse des cartes de lignes de champ).

Le potentiel électrostatique est positif au voisinage d'une charge positive et négatif au voisinage d'une charge négative. Son intensité (en valeur absolue) est d'autant plus grande que la charge est importante.

On constate enfin que le potentiel décroît comme la norme du champ quand on s'éloigne des charges.

A. Applications directes du cours

1. Distribution de charge uniforme dans tout l'espace

Soit une distribution volumique de charges ρ continue et uniforme dans tout l'espace.

1. Que peut-on dire du champ électrostatique en un point M quelconque de l'espace ?
2. Le théorème de Gauss est-il valide dans ce cas ? Expliquer.

2. Propriétés des équipotentielles

Soient les surfaces équipotentielles correspondant à une distribution donnée de charges.

1. En tout point d'une même équipotentielle, le module du champ électrostatique est-il le même ?
2. En un point d'une équipotentielle, quelle est l'orientation du champ électrostatique ?
3. Une équipotentielle peut-elle se recouper ?
4. Deux surfaces équipotentielles peuvent-elles se couper ?

3. Lignes de champ

On donne les lignes de champ suivantes :

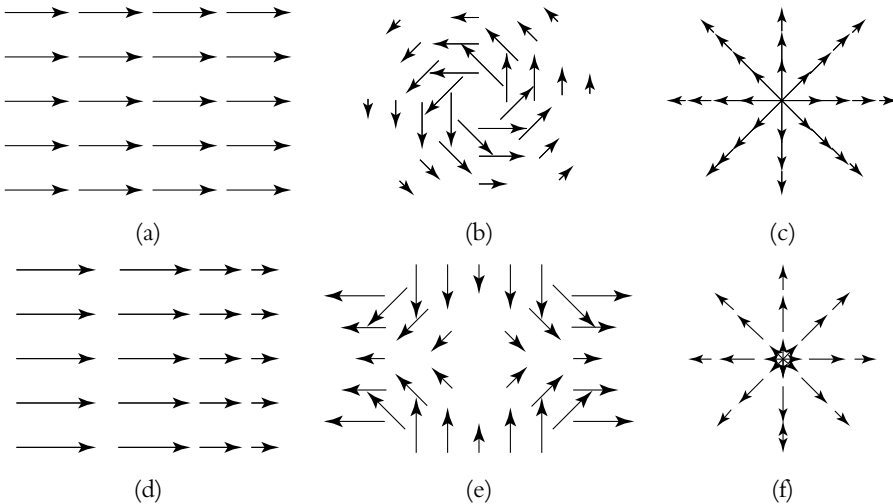


Figure 44.16

Préciser celles qui peuvent correspondre aux lignes de champ d'un champ électrostatique.

4. Flux du champ électrostatique à travers un cube

Dans la région de l'espace considérée règne un champ électrostatique $\vec{E}(M)$.

Soit un cube de côté a représenté sur la figure ci-contre :

Calculer le flux du champ électrostatique à travers le cube et en déduire la charge totale dans le cube si :

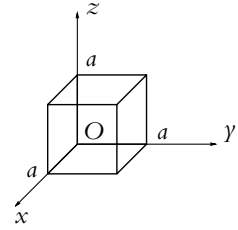


Figure 44.17

1. $\vec{E}(M)$ est uniforme,
2. $\vec{E}(M) = Cx\vec{u}_x$,
3. $\vec{E}(M) = Cx^2\vec{u}_x$,
4. $\vec{E}(M) = C(y\vec{u}_x + x\vec{u}_y)$.

B. Exercices et problèmes

1. Quatre charges aux sommets d'un carré

On considère la distribution constituée de quatre charges ponctuelles q de même valeur placées aux sommets d'un carré de côté a . On note O le centre du carré et on s'intéresse au potentiel et au champ électrostatique en un point M situé sur la droite (Ox) perpendiculaire au plan du carré et passant par O .

1. Déterminer le potentiel électrostatique créé en M par la distribution.
2. En déduire l'expression du champ électrostatique créé en M .
3. Trouver, en utilisant les symétries, la direction du champ électrostatique en M .
4. En exploitant les résultats de la question précédente, calculer directement le champ électrostatique créé en M .
5. Quelle est la parité du champ et du potentiel électrostatiques en x ?
6. Tracer la courbe donnant le potentiel électrostatique créé en M en fonction de x la distance de M au plan du carré.
7. Tracer la courbe donnant le champ électrostatique créé en M en fonction de x la distance de M au plan du carré.
8. Tracer l'allure des lignes de champs et de la trace des équipotentielle dans le plan Oyz .

2. Modèle d'atome d'hydrogène

On considère une distribution à symétrie sphérique créant le potentiel électrostatique suivant :

$$V(r) = \frac{q}{4\pi\epsilon_0} \frac{\exp\left(-\frac{r}{a}\right)}{r}.$$

1. Déterminer le champ électrostatique créé par cette distribution.
2. Étudier les cas limites $r \ll a$ et $r \gg a$.
3. Calculer le flux du champ électrostatique à travers une sphère de centre O et de rayon r . En déduire la charge $q(r)$ se trouvant à l'intérieur d'une sphère de rayon r centrée en O .

4. En déduire en O la présence d'une charge ponctuelle à déterminer. Quelle est la charge contenue dans tout l'espace ? Expliquer en quoi ce potentiel peut modéliser la distribution de charges d'un atome d'hydrogène.
5. Établir l'expression de la densité volumique de charge $\rho(r)$ à la distance r du point O .
6. On définit la densité de probabilité de présence $P(r)$ de l'électron dans le volume élémentaire $d\tau$ par $dq(r) = -eP(r) d\tau$ où $dq(r)$ est la charge contenue dans le volume $d\tau$. Exprimer $P(r)$ puis la probabilité $p(r)$ de trouver l'électron à une distance comprise entre r et $r + dr$ du noyau. Étudier la fonction $p(r)$. Comment peut-on interpréter a ? Donner l'ordre de grandeur de a .

3. Énergie de constitution d'une sphère chargée uniformément

Une sphère de rayon R porte la charge Q uniformément répartie en volume. On définit l'énergie de constitution (ou énergie coulombienne) de cette sphère comme le travail qu'il faut fournir pour la construire en prenant les charges à l'infini. On admet que cette énergie ne dépend pas de la façon dont on construit la sphère : on la construit par couches sphériques successives.

1. La sphère a un rayon $r < R$. Calculer le travail qu'il faut fournir pour augmenter le rayon de cette sphère de dr , en amenant les charges de l'infini.
2. En déduire l'expression de l'énergie de constitution de la sphère, en fonction de Q et de R .
3. Par analogie, en déduire l'énergie gravitationnelle d'une sphère de rayon R et de masse M .

4. Champ créé par une spire au voisinage de l'axe, d'après ENSTIM 2002

On donne une spire circulaire de rayon R , de centre O , d'axe Oz (figure 44.18). Cette spire porte une charge positive Q répartie uniformément avec densité linéique de charge λ en C.m^{-1} .

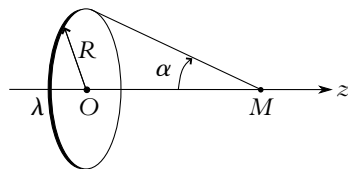


Figure 44.18

1. a) Montrer par des arguments de symétrie que, sur l'axe, le champ électrostatique $\vec{E}$ est porté par l'axe et prend la forme de $\vec{E} = E\vec{u}_z$ où $\vec{u}_z$ est un vecteur unitaire porté par l'axe Oz .
2. a) Comparer $E(z)$ et $E(-z)$.
- b) Calculer le champ électrostatique créé en un point M de l'axe tel que $OM = z$. On donnera le résultat en fonction de Q , la charge totale, du rayon R , de la permittivité du vide ϵ_0 et de la distance z .
- c) Tracer le graphe de la fonction $E(z)$.
3. On s'intéresse maintenant au champ électrostatique au voisinage de l'axe. On calcule donc le champ en un point M défini par des coordonnées cylindriques (r, θ, z) .
- a) Montrer par des arguments de symétrie très précis, qu'en M , le champ $\vec{E}$ n'a pas de composante orthoradiale E_θ .
- b) Montrer que la norme de E ne dépend que de r et z .

c) Montrer qu'au voisinage de l'axe, le flux du champ $\vec{E}$ est conservatif. Que peut-on dire de sa circulation sur un contour fermé ?

d) Calculer le flux de $\vec{E}$ à travers une surface fermée cylindrique d'axe Oz dont les bases sont des disques de rayon r petit et de cotes z et $z + dz$ (figure 44.19).

En déduire :

$$E_r(r, z) = -\frac{r}{2} \frac{dE_z(0, z)}{dz}$$

Établir l'expression de $E_r(r, z)$.

e) Calculer de même la circulation du champ électrostatique le long du petit rectangle de sommets $M(r, z)$, $N(r, z + dz)$, $P(r + dr, z + dz)$ et $Q(r + dr, z)$, en supposant r petit et dr et dz encore plus petits. En déduire que :

$$\frac{\partial E_z}{\partial r} = \frac{\partial E_r}{\partial z}$$

puis que :

$$E_z(r, z) = E_z(0, z) - \frac{r^2}{4} \frac{d^2 E_z}{dz^2}(0, z)$$

Établir l'expression de $E_z(r, z)$.

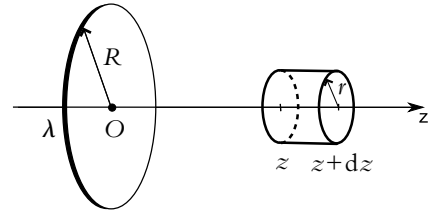


Figure 44.19

5. Champ électrostatique d'une distribution à symétrie sphérique, d'après CCP DEUG 2003

Soit le champ électrostatique créé dans le vide par une distribution de charges à symétrie sphérique. La composante radiale de ce champ dépend de la distance r au centre O de la distribution selon :

$$E_r = \begin{cases} \frac{k}{2\epsilon_0} & \text{si } 0 < r < R \\ \frac{kR^2}{2\epsilon_0 r^2} & \text{si } r > R \end{cases}$$

où k et R sont des constantes positives. On admet que dans les conditions du problème, la densité volumique de charges ρ vérifie

$$\rho = \frac{\epsilon_0}{r^2} \frac{d(r^2 E)}{dr}$$

1. De quelle(s) variable(s) dépend le champ électrostatique en coordonnées sphériques ?
2. Quelle est sa direction ?
3. Déterminer l'expression du potentiel V en supposant qu'il n'y a pas de charges à l'infini et que des charges infiniment éloignées n'interagissent pas entre elles.
4. Tracer l'allure du potentiel en fonction de r .
5. Exprimer la densité volumique de charges ρ .

6. Tracer l'allure de la densité volumique de charges en fonction de r .
7. Exprimer la charge dq d'une couche sphérique de centre O comprise entre r et $r + dr$.
8. En déduire l'expression de la charge totale q_0 de la distribution en fonction de k et R .
9. Etablir que pour $r > R$, on a l'équivalence de cette distribution avec une charge électrique ponctuelle dont on donnera la valeur et la position.
10. On considère maintenant deux charges ponctuelles q_0 et $q = -q_0$ placées à une distance a l'une de l'autre en O et A respectivement. On notera $\vec{u} = \frac{\vec{OA}}{OA}$.

Donner l'expression de la force $\vec{f}$ exercée par la charge q_0 sur la charge q .

11. Soit P un point équidistant de O et A . Déterminer la direction du champ électrostatique créé en P par les charges q et q_0 .
12. Préciser, à l'aide d'un schéma, son sens.
13. Si $OP = AP = a$, donner la valeur de la norme de ce champ.
14. Calculer le potentiel V' créé en P par les charges q et q_0 dans les mêmes conditions qu'à la question précédente.

Exemples de champs et potentiels électrostatiques

45

On présente dans ce chapitre quelques distributions classiques. C'est l'occasion de préciser la méthode à adopter pour calculer le champ et le potentiel électrostatiques créés par une distribution et d'analyser les résultats obtenus.

1. Méthodes de calcul

1.1 Démarche générale

Le calcul d'un champ électrostatique s'obtient en effectuant successivement les étapes suivantes¹ :

- Étude des invariances :

On a vu au cours des chapitres précédents que l'analyse des invariances permet de choisir le système de coordonnées adapté au problème :

- en cas d'invariance par translation, utilisation de coordonnées privilégiant un axe donc des coordonnées cartésiennes ou cylindriques,
- en cas d'invariance par rotation, utilisation de coordonnées précisant un angle donc des coordonnées cylindriques ou sphériques.

D'autre part, l'invariance par rapport à une ou plusieurs coordonnées réduit le nombre de variables d'espace à considérer : on a indépendance par rapport aux coordonnées selon lesquelles le système est invariant.

- Analyse des symétries :

Pour obtenir l'orientation du champ en un point, il faut rechercher les plans de symétrie et/ou d'antisymétrie passant par ce point et utiliser le fait que le champ

¹ On notera que l'ordre des deux premières étapes est interchangeable.

électrique appartient aux plans de symétrie et est perpendiculaire aux plans d'anti-symétrie. Il faut également tenir compte des propriétés suivantes :

- le champ électrostatique en un point M' symétrique d'un point M par rapport à un plan de symétrie est égal au symétrique par rapport au plan de symétrie du champ électrostatique en M ,
- le champ électrostatique en un point M' symétrique d'un point M par rapport à un plan d'antisymétrie est égal à l'opposé du symétrique par rapport au plan de symétrie du champ électrostatique en M .
- Choix d'une méthode de calcul parmi les trois exposées dans la suite.

1.2 Calcul direct de l'intégrale vectorielle

La première méthode est le calcul direct du champ à partir de sa définition : on calcule dans une base fixe (c'est-à-dire indépendante du point P) les trois composantes, ce qui revient :

- dans le cas d'une distribution continue à calculer trois intégrales² :

$$\vec{E}(M) = \frac{1}{4\pi\epsilon_0} \iiint_{P \in V} \frac{\rho(P)}{(PM)^3} \vec{PM} d\tau_P$$

l'intégrale portant sur l'ensemble des points P où se trouvent les charges,

- dans le cas d'une distribution discrète à calculer trois sommes discrètes portant sur l'ensemble des points P_i où se trouvent les charges :

$$\vec{E}(M) = \frac{1}{4\pi\epsilon_0} \sum_i \frac{q_i}{(P_i M)^3} \vec{P_i M}$$

Les invariances et les symétries permettent ici de simplifier les calculs. Certaines composantes sont nulles par symétrie : il sera inutile d'effectuer le calcul (au risque d'obtenir un résultat faux par erreur de calcul...).

1.3 Passage par le potentiel

Une deuxième possibilité consiste à passer par le potentiel et à calculer alors pour les distributions continues une seule intégrale scalaire³ :

$$V(M) = \frac{1}{4\pi\epsilon_0} \iiint_{P \in V} \frac{\rho(P)}{PM} d\tau_P$$

² Il s'agit d'intégrales simples, doubles ou triples suivant que la distribution est respectivement linéique, surfacique ou volumique. Sauf cas particulier, toutes les expressions pour les distributions continues seront données dans ce paragraphe dans le seul cas d'une distribution volumique.

³ On ne donne ici que l'expression pour une distribution volumique, celle d'une distribution surfacique ou linéique s'obtiennent par analogie.

ou pour les distributions discrètes une seule somme discrète :

$$V(M) = \frac{1}{4\pi\epsilon_0} \sum_i \frac{q_i}{P_i M}$$

Le champ électrostatique s'obtient alors en appliquant la relation :

$$\vec{E} = -\overrightarrow{\text{grad}} V$$

L'avantage de cette solution est de n'avoir qu'une seule intégrale ou somme à calculer au lieu de trois dans l'approche précédente. **Il faudra cependant faire attention à calculer l'expression du potentiel en tout point pour pouvoir ensuite calculer le gradient.** Tout comme la seule connaissance de la valeur d'une fonction en un point $f(x_0)$ ne permet pas de calculer la dérivée $f'(x_0)$ en ce point, la connaissance du potentiel en un point est insuffisante pour estimer le gradient de ce dernier et en déduire la valeur du champ.

On devra également tenir compte des propriétés du potentiel électrostatique :

- continuité du potentiel électrostatique pour un volume chargé ou une surface chargée mais discontinuité du potentiel électrostatique pour un fil chargé,
- le potentiel n'admet pas d'extrema en dehors des charges.

1.4 Passage par le théorème de Gauss

Lorsque l'analyse des symétries et des invariances a permis de déterminer la direction du champ, il ne reste plus qu'à estimer sa norme. Le théorème de Gauss fournit alors une réponse adaptée au problème. Pour pouvoir appliquer ce théorème, il convient dans un premier temps de rechercher les invariances et les symétries de la distribution. Cette méthode nécessite deux étapes supplémentaires.

- Détermination de la surface de Gauss : celle-ci est définie de manière à rendre élémentaire le calcul du flux. On choisira donc une surface composée de surfaces soit tangentes soit perpendiculaires au champ et sur lesquelles la norme du champ est uniforme. Dans le premier cas, le flux à travers la surface sera nul (le vecteur surface orientée et le champ électrique sont perpendiculaires : leur produit scalaire est nul). Dans le second, le produit scalaire entre vecteur surface orientée et champ électrique ne fait pas intervenir d'angle : les deux vecteurs sont colinéaires. De plus, puisque la norme du champ est uniforme sur cette surface, le calcul de l'intégrale sera très simple. Il est nécessaire de bien tenir compte du fait suivant : **la surface de Gauss doit passer par le point où on cherche à calculer le champ.**
- Calcul : il consiste à appliquer le théorème de Gauss compte tenu des résultats antérieurs.

D'autres méthodes basées sur les propriétés locales du champ électrostatique seront vues dans le cours de deuxième année.

2. Champ et potentiel créés par une sphère uniformément chargée en volume

Soit une sphère de centre O et de rayon R uniformément chargée en volume avec une densité volumique de charges ρ .

2.1 Analyse des invariances

- La distribution de charges présente une symétrie sphérique : on choisit donc les coordonnées sphériques.
- La distribution de charges est invariante par toute rotation autour du centre O : la norme du champ électrostatique $||\vec{E}(M)||$ et le potentiel électrostatique $V(M)$ sont donc indépendants des coordonnées angulaires θ et φ . Seule compte la distance r du point où on cherche le champ au centre de la sphère :

$$||\vec{E}(M)|| = E(r) \quad \text{et} \quad V(M) = V(r)$$

2.2 Analyse des symétries

Tout plan passant par le centre de la sphère et le point M est plan de symétrie des sources : le champ électrostatique appartient à tous ces plans et on en déduit la nature radiale de ce dernier

$$\vec{E}(M) = E(r)\vec{u}_r$$

Il ne reste plus que la norme du champ à déterminer : on va appliquer le théorème de Gauss.

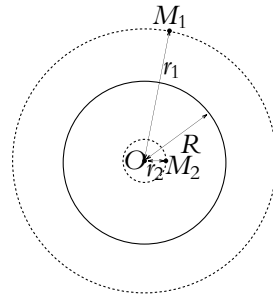


Figure 45.1 Sphère uniformément chargée.

2.3 Choix de la surface de Gauss

Le champ étant radial, on choisit une surface perpendiculaire au champ de façon à ce que le champ électrique et le vecteur surface orientée soient colinéaires en tout point. La surface de Gauss est donc la sphère de centre O (centre de la sphère chargée) et de rayon r . On notera qu'il s'agit d'une équipotentielle et que le module du champ est le même en tout point de la sphère.

2.4 Calcul

Le flux s'écrit :

$$\Phi = \oint_{M \in S} \vec{E}(M) \cdot d\vec{S}_M = \oint_{M \in S} E(r) dS_M$$

puisque $\vec{E}(M)$ est colinéaire à $d\vec{S}_M$. Finalement

$$\Phi = E(r) \oint \iint \iint_{M \in S} dS_M = 4\pi r^2 E(r)$$

car $E(r)$ est constant sur la surface $\mathcal{S}$ et que la surface de la sphère est $4\pi r^2$.

L'application du théorème de Gauss nécessite de calculer la charge intérieure à la surface de Gauss considérée. Il convient de distinguer deux cas :

- M est à l'intérieur de la sphère (point M_2 de la figure 4.5.1, avec $r = r_2$) : $r < R$ et $Q_{\text{int}} = \frac{4}{3}\pi r^3 \rho$ d'où :

$$4\pi r^2 E(r) = \frac{4\pi r^3 \rho}{3\epsilon_0}$$

et :

$$E(r) = \frac{\rho r}{3\epsilon_0} = \frac{Q_{\text{tot}}}{4\pi\epsilon_0 R^3} r$$

puisque la charge totale de la sphère vaut : $Q_{\text{tot}} = \frac{4}{3}\pi R^3 \rho$.

- M est à l'extérieur de la sphère (point M_1 de la figure 4.5.1, avec $r = r_1$) : $r > R$ et $Q_{\text{int}} = \frac{4}{3}\pi R^3 \rho$ d'où :

$$4\pi r^2 E(r) = \frac{4\pi R^3 \rho}{3\epsilon_0}$$

et :

$$E(r) = \frac{\rho R^3}{3\epsilon_0 r^2} = \frac{Q_{\text{tot}}}{4\pi\epsilon_0 r^2} \frac{1}{r^2}$$

en tenant compte de la valeur de la charge totale. On peut remarquer qu'à l'extérieur de la sphère, le champ est identique à celui créé par une charge ponctuelle placée au centre de la sphère et dont la valeur est la charge totale de la sphère.

On remarque que le module du champ est continu en $r = R$.

2.5 Détermination du potentiel

On en déduit le potentiel électrostatique par la relation : $\vec{E} = -\vec{\text{grad}}V$ soit :

$$-\frac{dV}{dr} = E(r)$$

Le potentiel ne dépend que de r , on doit donc considérer la dérivée d'une fonction d'une seule variable et ne pas recourir aux dérivées partielles.

On distingue à nouveau les deux cas :

- M est à l'intérieur de la sphère : $r < R$

$$-\frac{dV}{dr} = \frac{Q_{\text{tot}}}{4\pi\epsilon_0 R^3} r$$

d'où :

$$V = -\frac{Q_{\text{tot}}}{4\pi\epsilon_0} \frac{r^2}{2R^3} + C$$

où C est une constante ;

- M est à l'extérieur de la sphère : $r > R$

$$-\frac{dV}{dr} = \frac{Q_{\text{tot}}}{4\pi\epsilon_0} \frac{1}{r^2}$$

d'où :

$$V = \frac{Q_{\text{tot}}}{4\pi\epsilon_0} \frac{1}{r} + C'$$

où C' est une autre constante.

Les deux constantes C et C' se déterminent en utilisant :

- la continuité du potentiel à la surface de la sphère ($r = R$),
- les conditions aux limites en choisissant classiquement le potentiel nul à l'infini :

$$\lim_{r \rightarrow \infty} V(r) = 0$$

On en déduit que $C' = 0$ et

$$V(r > R) = \frac{Q_{\text{tot}}}{4\pi\epsilon_0} \frac{1}{r}$$

ainsi que

$$V(R) = \frac{Q_{\text{tot}}}{4\pi\epsilon_0} \frac{1}{R} = -\frac{Q_{\text{tot}}}{4\pi\epsilon_0} \frac{1}{2R} + C$$

d'où

$$C = \frac{3Q_{\text{tot}}}{8\pi\epsilon_0 R}$$

et

$$V(r < R) = \frac{Q_{\text{tot}}}{4\pi\epsilon_0 R} \left(\frac{3}{2} - \frac{r^2}{2R^2} \right)$$

D'où les allures des fonctions $E(r)$ et $V(r)$:

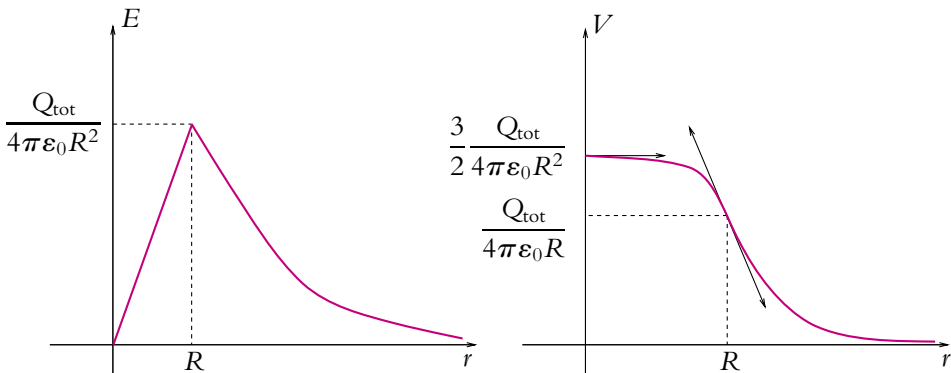


Figure 45.2 Norme du champ et potentiel électrostatiques créés par une sphère uniformément chargée.

On peut formuler les remarques suivantes :

- Le champ et le potentiel électrostatiques à l'extérieur de la sphère sont les mêmes que ceux d'une charge ponctuelle dont la valeur est la charge totale de la sphère et qui serait placée au centre de la sphère. Cela lève le problème de discontinuité rencontré pour une charge ponctuelle : on considère une certaine extension spatiale pour s'affranchir des discontinuités. On reverra ce point dans le paragraphe 4.7 du chapitre, qui s'intéresse à ce problème sur l'exemple du plan infini.
- Le potentiel ne présente pas de point anguleux en $r = R$ du fait que le champ électrostatique est continu en ce point-là. En effet, la dérivée du potentiel est ici égale au champ électrostatique, elle est continue en ce point.

2.6 Analogie avec le champ gravitationnel

On s'intéresse au champ gravitationnel créé par une sphère de centre O et de rayon R de masse volumique uniforme $\rho = \frac{M_{\text{tot}}}{\frac{4}{3}\pi R^3}$.

a) Analyse des invariances

La distribution de masse présente une symétrie sphérique : on utilise donc les coordonnées sphériques.

Elle est invariante par toute rotation autour de O : la norme du champ gravitationnel $\|\vec{A}(M)\|$ en un point M ne dépend que de la distance r de M à O : $\|\vec{A}(M)\| = A(r)$.

b) Analyse des symétries

Tout plan passant par O et M est plan de symétrie des sources : le champ gravitationnel $\vec{A}(M)$ appartient donc à tous ces plans. On en déduit que le champ est radial : $\vec{A}(M) = A(r)\vec{u}_r$. Il ne reste plus que le module du champ à déterminer : on applique donc le théorème de Gauss.

c) Choix de la surface de Gauss

Le champ étant radial, on choisit une surface perpendiculaire au champ de façon à ce que le champ électrique et le vecteur surface orientée soient colinéaires en tout point. La surface de Gauss est donc la sphère de centre O et de rayon r . On notera que le module du champ est le même en tout point de la sphère.

d) Calcul du module du champ

Le flux s'écrit :

$$\Phi = \oint_{M \in S} \vec{A}(M) \cdot d\vec{S}_M = \oint_{M \in S} A(r) dS_M = A(r) \oint_{M \in S} dS_M = 4\pi r^2 A(r)$$

L'application du théorème de Gauss nécessite de calculer la masse intérieure à la surface de Gauss considérée. Il convient de distinguer deux cas :

- M est à l'intérieur de la sphère : $r < R$ et $M_{\text{int}} = \frac{4}{3}\pi r^3 \rho$ d'où

$$4\pi r^2 A(r) = -4\pi G \frac{4}{3}\pi r^3 \rho$$

et

$$A(r) = -\frac{4}{3}\pi G \rho r = -GM_{\text{tot}} \frac{r}{R^3}$$

- M est à l'extérieur de la sphère : $r > R$ et $M_{\text{int}} = \frac{4}{3}\pi R^3 \rho$ d'où

$$4\pi r^2 A(r) = -4\pi G \frac{4}{3}\pi R^3 \rho$$

et

$$A(r) = -\frac{4}{3}\pi G R^3 \rho = -\frac{GM_{\text{tot}}}{r^2}$$

On peut remarquer qu'à l'extérieur de la sphère, le champ est identique à celui créé par une masse ponctuelle placée au centre de la sphère et dont la valeur est la masse totale de la sphère. On a utilisé ce résultat lors de l'analyse qualitative du phénomène de marée dans le cours de mécanique.

On remarque que le module du champ est continu en $r = R$.

e) Analogie avec le champ électrostatique

Tout ce qui vient d'être fait dans ce paragraphe aurait pu être obtenu directement en utilisant l'analogie.

Les résultats concernant les invariances et les symétries sont immédiates du fait d'un comportement analogue des champs électrostatique et gravitationnel pour une même situation, ici celle d'une sphère à distribution uniforme.

Quant à la valeur du champ, il suffit de remplacer formellement ϵ_0 par $-\frac{1}{4\pi G}$ et Q_{tot} par M_{tot} pour obtenir le résultat.

Le champ électrostatique est orienté selon $\vec{u}_r$ si $\rho > 0$ et selon $-\vec{u}_r$ sinon ; le champ gravitationnel est toujours orienté selon $-\vec{u}_r$.

3. Champ et potentiel créés par un cylindre ou un fil uniformément chargé

Soit un cylindre d'axe Oz et de base un disque de centre O et de rayon a uniformément chargé en volume avec une densité volumique de charges ρ . Le cas du fil correspond à faire tendre a vers 0 et à considérer une distribution linéique de densité λ .

3.1 Analyse des invariances

- On suppose que la distance r du point M où on cherche à calculer le champ électrostatique est faible devant la hauteur du cylindre de manière à pouvoir le considérer comme infini. Avec cette hypothèse, la distribution de charges est invariante par translation le long de Oz : la norme du champ et le potentiel électrostatiques ne dépendent donc pas de z .
- La distribution est également invariante par rotation autour de l'axe Oz . En coordonnées cylindriques, la norme du champ et le potentiel électrostatique ne dépendent donc pas de θ .

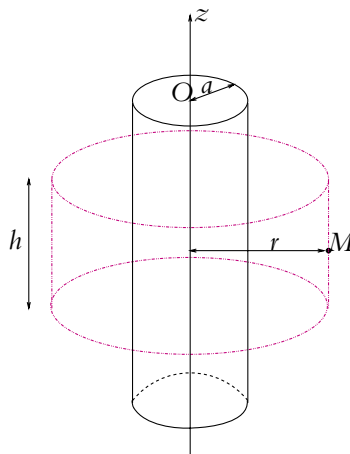


Figure 45.3 Cylindre uniformément chargé.

D'après les deux points précédents, la norme du champ et le potentiel électrostatiques ne dépendent que de la distance r du point M où on cherche le champ à l'axe Oz .

$$||\vec{E}(M)|| = E(r) \quad \text{et} \quad V(M) = V(r)$$

3.2 Analyse des symétries

Le plan contenant le point M et l'axe Oz ainsi que le plan contenant le point M et perpendiculaire à l'axe Oz sont plans de symétrie des sources. Le champ appartient à ces plans : $\vec{E}(M)$ n'a de composantes ni suivant $\vec{u}_z$ ni suivant $\vec{u}_\theta$ donc

$$\vec{E}(M) = E(r)\vec{u}_r$$

Comme pour le cas de la sphère, on va donc appliquer le théorème de Gauss.

3.3 Choix de la surface de Gauss

On choisit comme surface de Gauss un cylindre de hauteur h centré sur l'axe Oz et passant par M .

En effet, le flux est nul sur les deux bases du cylindre : le champ est, compte tenu des symétries, colinéaire à ces surfaces donc perpendiculaire aux vecteurs surfaces orientées.

Il ne subsiste donc dans le calcul du flux que le flux latéral, ce dernier s'obtiendra facilement puisque le champ est alors colinéaire au vecteur surface orientée et que sa norme est constante sur la surface latérale.

Le flux de $\vec{E}$ à travers le cylindre choisi vaut donc :

$$\Phi = 2\pi rhE(r)$$

3.4 Calcul pour le fil

Le calcul de la charge intérieure à la surface de Gauss donne $Q_{\text{int}} = \lambda h$. On en déduit $2\pi rhE(r) = \frac{\lambda h}{\epsilon_0}$ et on retrouve le résultat déjà obtenu par le calcul direct au chapitre sur le champ électrostatique :

$$E(r) = \frac{\lambda}{2\pi\epsilon_0 r}$$

3.5 Calcul pour le cylindre

Le calcul de la charge intérieure à la surface de Gauss nécessite de distinguer deux cas :

- M à l'intérieur du cylindre $r < a$: $Q_{\text{int}} = \pi r^2 h \rho$ d'où :

$$2\pi rhE(r) = \frac{\rho\pi r^2 h}{\epsilon_0}$$

soit :

$$E(r) = \frac{\rho r}{2\epsilon_0}$$

- M à l'extérieur du cylindre $r > a$: $Q_{\text{int}} = \pi a^2 h \rho$ d'où

$$2\pi rhE(r) = \frac{\rho\pi a^2 h}{\epsilon_0}$$

soit :

$$E(r) = \frac{\rho a^2}{2\epsilon_0 r}$$

On notera qu'ici aussi le champ électrostatique est continu en $r = a$.

En outre, la grandeur h , qui est choisie arbitrairement, s'élimine au cours du calcul ce qui justifie *a posteriori* son utilisation.

3.6 Détermination du potentiel pour le cylindre

On en déduit le potentiel par la relation :

$$-\frac{dV}{dr} = E(r)$$

soit :

- si M est à l'intérieur du cylindre $r < a$: $-\frac{dV}{dr} = \frac{\rho r}{2\epsilon_0}$ et :

$$V = -\frac{\rho r^2}{4\epsilon_0} + C$$

où C est une constante ;

- si M est à l'extérieur du cylindre $r > a$:

$$-\frac{dV}{dr} = \frac{\rho a^2}{2\varepsilon_0 r}$$

et

$$V = -\frac{\rho a^2}{2\varepsilon_0} \ln \frac{r}{r_0} + C'$$

où r_0 et C' sont des constantes.

Les constantes C et C' sont déterminées par les conditions aux limites et la continuité du potentiel en $r = a$.

Ici on a des charges à l'infini : il est impossible de choisir un potentiel nul à l'infini. On fixe donc le potentiel nul pour une distance r_0 de l'axe : $V(r_0) = 0$. Donc :

$$V(r > a) = -\frac{\rho a^2}{2\varepsilon_0} \ln \frac{r}{r_0}$$

La deuxième constante s'obtient par continuité en $r = a$:

$$-\frac{\rho a^2}{2\varepsilon_0} \ln \frac{a}{r_0} = -\frac{\rho a^2}{4\varepsilon_0} + C$$

donc :

$$V(r < a) = -\frac{\rho a^2}{2\varepsilon_0} \ln \frac{a}{r_0} - \frac{\rho}{4\varepsilon_0} (r^2 - a^2)$$

D'où les allures des fonctions $E(r)$ et $V(r)$:

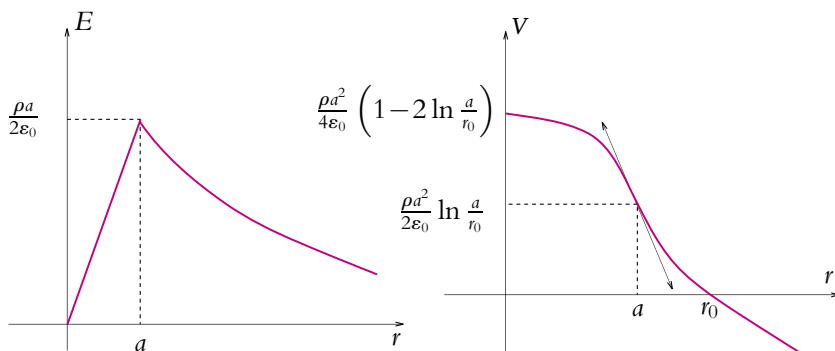


Figure 45.4 Norme du champ et potentiel électrostatique d'un cylindre uniformément chargé.

On note que le champ électrostatique étant continu en $r = a$, le potentiel électrostatique ne présente pas de point anguleux en ce point.

4. Champ et potentiel créés par un plan uniformément chargé

Soit un plan $\mathcal{P}$ infini portant une charge surfacique σ uniforme sur tout le plan. On rappelle qu'un plan peut être considéré comme infini si la distance du point considéré au plan est petite devant les dimensions du plan.

4.1 Analyse des invariances

Les sources sont invariantes par toute translation de vecteur parallèle à $\mathcal{P}$: le champ ne dépend que de la distance au plan $\mathcal{P}$.

4.2 Analyse des symétries

- Tout plan perpendiculaire au plan $\mathcal{P}$ et passant par le point M est plan de symétrie des sources : le champ électrostatique appartient à ce plan et est donc perpendiculaire au plan $\mathcal{P}$ en tout point.
- Le plan $\mathcal{P}$ est un plan de symétrie des sources : les champs électrostatiques créés en deux points symétriques l'un de l'autre par rapport au plan $\mathcal{P}$ sont symétriques l'un de l'autre. On a donc $\vec{E}(M) = -\vec{E}(N)$ avec $N = \text{sym}_{\mathcal{P}}(M)$.

On notera que le champ électrostatique n'est pas défini en un point du plan $\mathcal{P}$. On verra dans la suite que le champ présente une discontinuité à la traversée du plan.

La seule information manquante est la norme du champ pour un point M distant de $z \neq 0$ du plan $\mathcal{P}$: on va donc appliquer le théorème de Gauss.

4.3 Détermination de la surface de Gauss

La surface de Gauss est la surface fermée indiquée sur la figure ci-contre. Elle est formée par deux éléments de surface ΔS_1 et ΔS_2 parallèles au plan $\mathcal{P}$ et distante de z du plan $\mathcal{P}$ de part et d'autre de ce dernier et un tube de champ s'appuyant sur une courbe fermée appartenant à $\mathcal{P}$ et de surface ΔS . On choisit les surfaces ΔS_1 et ΔS_2 à une distance z du plan pour déterminer la valeur du champ à cette distance du plan et on utilise la propriété :

$$E(z) = -E(-z)$$

Le champ est colinéaire aux surfaces orientées pour ΔS_1 et ΔS_2 , de norme constante sur ces deux surfaces, et perpendiculaire aux surfaces orientées sur le tube de champ, ce qui correspond aux deux cas utilisables pour appliquer le théorème de Gauss.

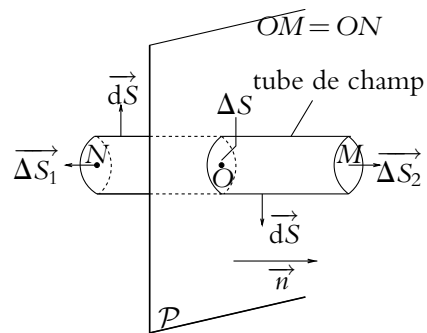


Figure 45.5 Plan uniformément chargé.

4.4 Calcul

Le flux du champ électrostatique $\vec{E}$ s'écrit :

$$\Phi = \Phi_{\text{latéral}} + \Phi_{\Delta S_1} + \Phi_{\Delta S_2}$$

$\Phi_{\text{latéral}} = 0$ car le champ électrostatique $\vec{E}$ perpendiculaire aux surfaces orientées élémentaires $d\vec{S}$ du tube de champ.

Compte tenu des symétries, $\Phi_{\Delta S_1} = \Phi_{\Delta S_2} = E(z)\Delta S$ car le champ électrostatique $\vec{E}$ est parallèle aux surfaces orientées ΔS_1 et ΔS_2 .

Donc

$$\Phi = 2E(z)\Delta S$$

et le théorème de Gauss s'écrit :

$$\Phi = \frac{Q_{\text{int}}}{\epsilon_0} = \frac{\sigma \Delta S}{\epsilon_0}$$

On en déduit :

$$\vec{E}(M) = \frac{\sigma}{2\epsilon_0} \vec{n}$$

en notant $\vec{n}$ le vecteur unitaire normal au plan $\mathcal{P}$ dirigé vers le point M où on calcule le champ.

On notera qu'on a choisi arbitrairement la surface ΔS et qu'elle s'élimine, ce qui justifie *a posteriori* ce choix arbitraire.

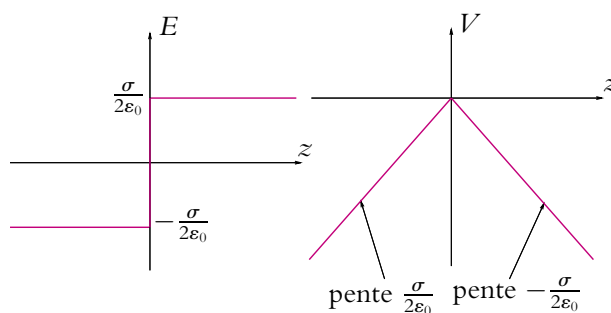


Figure 45.6 Norme du champ et potentiel électrostatique d'un plan uniformément chargé.

4.5 Détermination du potentiel

On peut alors en déduire le potentiel électrostatique par la relation : $\vec{E} = -\vec{\text{grad}} V$ soit :

$$-\frac{dV}{dz} = \frac{\sigma}{2\epsilon_0} \text{signe}(z) \quad \text{et} \quad V = -\frac{\sigma}{2\epsilon_0} |z| + \text{constante}$$

On choisit arbitrairement un potentiel nul sur le plan $V(0) = 0$ et

$$V = -\frac{\sigma}{2\epsilon_0}|z|$$

On peut remarquer la discontinuité du champ et simultanément le point anguleux du potentiel, ce qui est en parfait accord avec la relation existant entre champ et potentiel.

4.6 Application au condensateur plan

Dans ce paragraphe, on ne tient pas compte des effets de bord : cela revient à considérer que les armatures du condensateur plan sont des plans infinis et à négliger ce qui se passe sur les bords des plans.

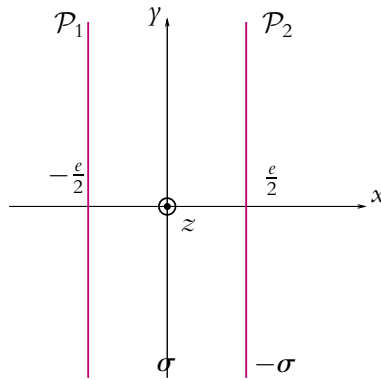


Figure 45.7 Condensateur plan.

Dans ces conditions, un condensateur plan est assimilable à l'association de deux plans infinis uniformément chargés, l'un de densité de charge $+\sigma$, l'autre de densité de charge $-\sigma$; ces deux plans sont distants de e . On peut exprimer σ en fonction de la charge totale Q et de la surface S d'une armature par :

$$\sigma = \frac{Q}{S}$$

On suppose que le plan de densité de charges $-\sigma$ est le plan $\mathcal{P}_1$ d'équation $x = -\frac{e}{2}$ et celui de densité de charges $+\sigma$ le plan $\mathcal{P}_2$ d'équation $x = \frac{e}{2}$. En utilisant les résultats obtenus aux paragraphes précédents, le champ électrostatique créé par $\mathcal{P}_1$ est :

$$\vec{E}_1 = \begin{cases} \frac{\sigma}{2\epsilon_0} \vec{u}_x & \text{si } x \leq -\frac{e}{2} \\ -\frac{\sigma}{2\epsilon_0} \vec{u}_x & \text{si } x \geq -\frac{e}{2} \end{cases}$$

et le champ électrostatique créé par $\mathcal{P}_2$ est :

$$\vec{E}_2 = \begin{cases} -\frac{\sigma}{2\epsilon_0} \vec{u}_x & \text{si } x \leq \frac{e}{2} \\ \frac{\sigma}{2\epsilon_0} \vec{u}_x & \text{si } x \geq \frac{e}{2} \end{cases}$$

Par le principe de superposition, le champ électrostatique créé par la distribution de charges du condensateur est :

$$\vec{E} = \begin{cases} \vec{0} & \text{si } |x| \geq \frac{e}{2} \\ -\frac{\sigma}{\epsilon_0} \vec{u}_x & \text{si } |x| \leq \frac{e}{2} \end{cases}$$

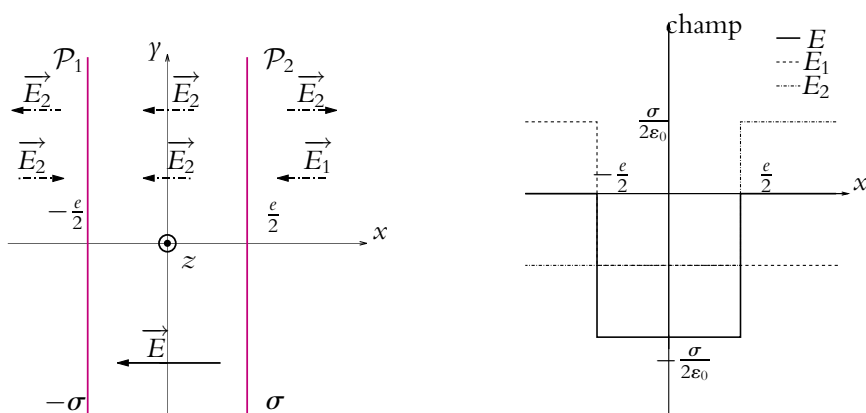


Figure 45.8 Champ électrostatique créé par la distribution d'un condensateur plan.

On en déduit le potentiel V de cette distribution en utilisant la relation :

$\vec{E} = -\text{grad}V$ soit ici :

$$\frac{dV}{dx} = -\left(-\frac{\sigma}{\epsilon_0}\right) = \frac{\sigma}{\epsilon_0}$$

et par intégration :

$$V = \frac{\sigma}{\epsilon_0}x + C$$

en notant C une constante pouvant être définie en choisissant une origine des potentiels.

À partir de l'expression du potentiel, on peut calculer la tension aux bornes du condensateur plan ; il s'agit de la différence de potentiel entre les deux armatures

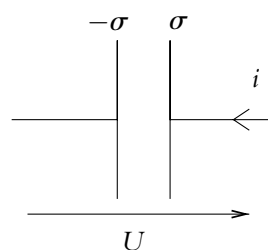


Figure 45.9 Représentation schématique du condensateur.

soit :

$$U = V\left(\frac{e}{2}\right) - V\left(-\frac{e}{2}\right) = \frac{\sigma e}{\epsilon_0}$$

Le choix du sens est lié aux conventions choisies en électrocinétique.

Or, comme $\sigma = \frac{Q}{S}$, on en déduit la relation fondamentale du condensateur (déjà utilisée en électrocinétique) entre la charge Q des armatures et la tension aux bornes du condensateur :

$$U = \frac{e}{S\epsilon_0} Q \quad \text{ou} \quad Q = \frac{S\epsilon_0}{e} U$$

On a donc bien la proportionnalité entre la charge des armatures et la tension aux bornes du condensateur. On déduit de la relation précédente la valeur de la capacité d'un condensateur plan en fonction de la surface S des armatures, de la distance e entre les armatures et de ϵ_0 la permittivité électrique dans le vide :

$$C = \frac{S\epsilon_0}{e}$$

On note que l'unité de ϵ_0 est F.m^{-1} par homogénéité.

► **Remarque :** dans le cas où le diélectrique n'est pas le vide, on doit remplacer la permittivité électrique dans le vide ϵ_0 par le produit $\epsilon_r \epsilon_0$ où ϵ_r est la permittivité électrique relative caractéristique du diélectrique. Ces aspects seront en partie développés dans le cours de seconde année PC/PC*.

4.7 Analyse des discontinuités

Dans le cas d'un plan infini uniformément chargé, le champ électrostatique est discontinu du fait de la traversée d'un plan chargé avec une densité surfacique de charges σ . Le but de ce paragraphe est de montrer qu'on peut faire disparaître cette discontinuité en modifiant le modèle.

Pour cela, on tient compte de l'épaisseur e du plan en considérant la distribution volumique de charges suivante :

$$\rho = \begin{cases} \rho_0 \text{ constante} & \text{si } |z| < \frac{e}{2} \\ 0 & \text{si } |z| > \frac{e}{2} \end{cases}$$

avec $\sigma S = \rho_0 e S$ pour avoir la même quantité de charges pour une même surface pour les deux distributions, soit :

$$\rho_0 = \frac{\sigma}{e}$$

Les invariances et les symétries sont les mêmes que pour le plan infini uniformément chargé : $\vec{E}(M) = E(z)\vec{u}_z$

On utilise la même surface de Gauss. Pour un point extérieur à la distribution volumique de charges, $Q_{\text{int}} = \rho_0 e \Delta S$ d'où

$$\Phi = 2E(z)\Delta S = \rho_0 e \Delta S$$

et on retrouve le résultat :

$$\vec{E}(M) = \frac{\sigma}{2\epsilon_0} \vec{n}$$

Pour un point à l'intérieur de la distribution volumique, la charge intérieure devient :

$$Q_{\text{int}} = 2\rho_0 z \Delta S$$

d'où

$$\Phi = 2E(z)\Delta S = 2\rho_0 z \Delta S = 2\frac{\sigma}{e} z \Delta S$$

et on obtient le résultat :

$$\vec{E}(M) = \frac{\sigma}{\epsilon_0} \frac{z}{e} \vec{n}$$

On peut en déduire le potentiel (le seul calcul à faire est celui à l'intérieur de la distribution) :

$$-\frac{dV}{dz} = \frac{\sigma z}{\epsilon_0 e}$$

donc :

$$V(M) = -\frac{\sigma}{2\epsilon_0 e} z^2 + C$$

On choisit la constante nulle en $z = 0$ et puisqu'à l'extérieur de la distribution, on a :

$$V = -\frac{\sigma}{2\epsilon_0} |z| + C'$$

la continuité de V en $z = \frac{e}{2}$ donne alors :

$$C' - \frac{\sigma e}{4\epsilon_0} = -\frac{\sigma e}{8\epsilon_0}$$

donc $C' = \frac{\sigma e}{8\epsilon_0}$ et

$$V = \begin{cases} -\frac{\sigma}{2\epsilon_0 e} z^2 & \text{si } |z| < \frac{e}{2} \\ -\frac{\sigma}{2\epsilon_0} |z| + \frac{\sigma e}{8\epsilon_0} & \text{si } |z| > \frac{e}{2} \end{cases}$$

D'où les allures des fonctions $E(z)$ et $V(z)$ données par la figure 45.10.

On remarquera qu'en faisant tendre e vers 0 tout en gardant une densité de charges σ constante, on retrouve le cas du plan qui présente une discontinuité.

Ce qui vient d'être fait sur l'exemple du plan pourrait être réalisé dans tous les cas de discontinuité rencontrés.

Dans tous les exemples où apparaît une discontinuité du champ électrostatique à la traversée d'une surface chargée avec une densité surfacique de charges σ , on peut remarquer que :

$$\vec{E}_2(M) - \vec{E}_1(M) = \frac{\sigma(M)}{\epsilon_0} \vec{n}_{12}$$

en notant $\vec{n}_{12}$ la normale dirigée de 1 vers 2. Ce résultat sera généralisé en deuxième année.

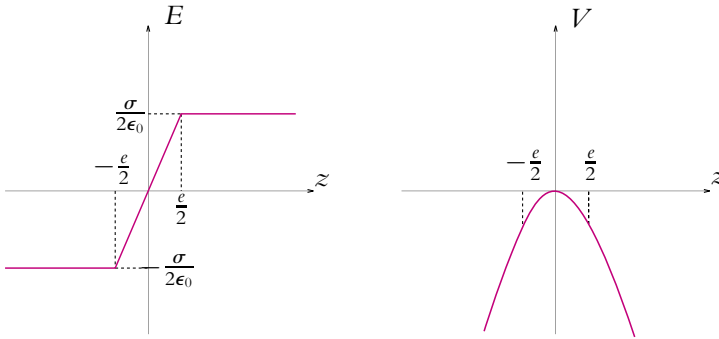


Figure 45.10 Norme du champ et potentiel électrostatique d'un plan uniformément chargé en tenant compte de son épaisseur

5. Champ et potentiel créés par un disque uniformément chargé le long de son axe

Soit un disque de centre O et de rayon R uniformément chargé en surface (avec une densité surfacique de charges σ) dont on veut déterminer le champ et le potentiel sur l'axe qui lui est perpendiculaire et qui passe par son centre.

5.1 Analyse des invariances

- On se place en coordonnées cylindriques du fait de la symétrie de la distribution de charges.
- On cherche le champ le long de l'axe Oz donc la seule composante non nulle de $\vec{M}$ est z . On a donc :

$$\vec{E}(M) = \vec{E}(z)$$

- Bien que ce soit inutile pour le problème considéré ici, on peut signaler que l'invariance de la distribution par rotation autour de l'axe Oz implique que le champ ne dépend pas de la composante angulaire θ et que pour un point en dehors de l'axe, on a :

$$||\vec{E}(M)|| = E(r, z)$$

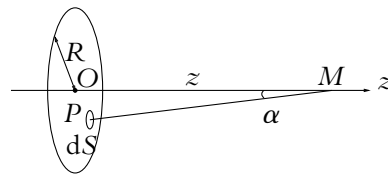


Figure 45.11 Disque uniformément chargé.

5.2 Analyse des symétries

Tout plan contenant l'axe Oz est plan de symétrie pour les sources : le champ appartient à ces plans. Pour un point de l'axe, le champ est colinéaire à l'axe.

Donc :

$$\vec{E}(M) = E_z(z)\vec{u}_z$$

En outre, le plan du disque étant un plan de symétrie de la distribution de charge, on a :

$$E_z(-z) = -E_z(z)$$

puisque le champ en un point repéré par $-z$ est le champ du point symétrique repéré par z par rapport au plan de symétrie des sources et est donc égal au symétrique du champ au point repéré par z .

5.3 Calcul direct

Bien que le seul point restant soit la norme du champ, on ne peut pas appliquer le théorème de Gauss. En dehors de l'axe (où on connaît l'orientation du champ), on ne sait rien de l'orientation du champ électrostatique $\vec{E}(M)$. De plus, le champ dépend alors de r car $r \neq 0$.

On va donc appliquer la méthode du calcul direct. Supposons que $z > 0$.

$$\vec{E}(M) = \frac{1}{4\pi\epsilon_0} \iint_{P \in \text{disque}} \sigma \frac{\vec{PM}}{(PM)^3} dS$$

soit en projection sur Oz qui est la seule composante non nulle d'après l'étude préalablement menée :

$$\begin{aligned} E_z &= \frac{1}{4\pi\epsilon_0} \iint_{P \in \text{disque}} \frac{\sigma \vec{PM} \cdot \vec{u}_z}{(PM)^3} r dr d\theta \\ &= \frac{1}{4\pi\epsilon_0} \iint_{P \in \text{disque}} \frac{\sigma PM \cos \alpha}{(PM)^3} r dr d\theta \\ &= \frac{1}{4\pi\epsilon_0} \int_0^{2\pi} d\theta \int_0^R \frac{\sigma z}{(z^2 + r^2)^{\frac{3}{2}}} r dr \\ &= \frac{\sigma z}{4\epsilon_0} \int_0^{R^2} \frac{dX}{(z^2 + X)^{\frac{3}{2}}} \\ &= \frac{\sigma}{2\epsilon_0} \left(1 - \frac{z}{\sqrt{z^2 + R^2}} \right) \end{aligned}$$

Pour z quelconque :

$$\vec{E}(M) = \frac{\sigma}{2\epsilon_0} \left(1 - \frac{|z|}{\sqrt{z^2 + R^2}} \right) \text{signe}(z) \vec{u}_z$$

5.4 Lien entre champ et potentiel

a) Obtention du potentiel par la relation champ - potentiel

On a la relation :

$$\frac{dV}{dz} = -E_z = -\frac{\sigma}{2\epsilon_0} \left(1 - \frac{z}{\sqrt{z^2 + R^2}} \right)$$

d'où par intégration :

$$V = -\frac{\sigma}{2\epsilon_0} \int \left(1 - \frac{u}{\sqrt{u^2 + R^2}} \right) du = -\frac{\sigma}{2\epsilon_0} \left(|z| - \sqrt{z^2 + R^2} \right)$$

en choisissant la constante de manière à avoir un potentiel nul à l'infini.

b) Calcul direct du potentiel

$$\begin{aligned} V &= \frac{1}{4\pi\epsilon_0} \iint_{P \in \text{disque}} \frac{\sigma dS}{PM} \\ &= \frac{\sigma}{4\pi\epsilon_0} \int_0^{2\pi} d\theta \int_0^R \frac{r dr}{\sqrt{z^2 + r^2}} \\ &= \frac{\sigma}{2\epsilon_0} \left(\sqrt{z^2 + R^2} - |z| \right) \end{aligned}$$

On pourrait également mener le calcul en utilisant l'angle α comme variable.

On peut alors calculer le champ électrostatique. En effet, on obtient le même résultat en dérivant par rapport à z et en prenant $r = 0$ et $\theta = 0$ qu'en prenant $r = 0$ et $\theta = 0$ puis en dérivant par rapport à z d'après les propriétés de la dérivée partielle.

5.5 Courbes représentatives et commentaires

Les allures des fonctions $E(z)$ et $V(z)$ sont donc celles de la figure 45.12.

On notera la discontinuité du champ électrostatique et le point anguleux correspondant du potentiel électrostatique.

Dans le cas où le rayon tend vers l'infini, on retrouve le cas du plan infini en prenant la limite de l'expression précédente.

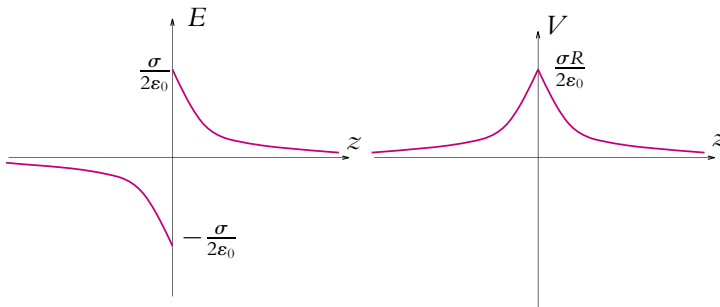


Figure 45.12 Norme du champ et potentiel électrostatique le long de l'axe d'un disque uniformément chargé.

A. Applications directes du cours

1. Champ sur l'axe d'un carré

- On considère un segment de fil, de longueur L , portant une densité linéique de charge λ . Calculer le champ électrique en un point M de la médiatrice, situé à une distance D du fil.
- On considère maintenant un carré formé de quatre segments semblables au précédent. On note Oz l'axe passant par le centre C du carré et perpendiculaire au plan xOy du carré. Déterminer le champ électrique en un point M de l'axe Oz situé à une distance z au-dessus du plan xOy .

2. Sphère portant une densité surfacique de charge non uniforme

On considère une sphère de centre O , d'axe Oz . Pour un point M de la sphère, on appelle θ l'angle que fait OM avec Oz . La sphère porte une densité surfacique de charge $\sigma(M) = \sigma_0 \cos \theta$, σ_0 étant une constante.

- Par l'étude des symétries de la distribution de charge, déterminer la direction du champ $\vec{E}$ en O .
- a) Déterminer la projection dE_z du champ élémentaire $d\vec{E}$ créé par une surface élémentaire dS .
b) Calculer le champ $\vec{E}$ en O .
- Déterminer le potentiel en O .

3. Les distributions de base avec le théorème de Gauss

En adoptant la démarche suivante, et sans regarder le cours :

- étude des symétries du champ électrostatique $\vec{E}$ pour en déduire la direction de $\vec{E}$,
- étude des invariances du champ électrostatique $\vec{E}$ pour en déduire les dépendances de $\|\vec{E}\|$,
- utilisation du théorème de Gauss pour déterminer complètement le champ $\vec{E}$ en tout point de l'espace,
- détermination de l'expression du potentiel électrostatique V ,
étudier les distributions suivantes :
 - un fil infini chargé avec la densité linéique uniforme λ .
 - un plan infini chargé avec la densité surfacique uniforme σ .
 - une boule de rayon R chargée avec la densité volumique uniforme ρ .

4. Boule avec cavité

On considère une boule de rayon R et de centre O uniformément chargée en volume (densité ρ). Cette boule possède une cavité vide de charge centrée en O' et de rayon a .

Par application du théorème de superposition, déterminer la champ électrostatique dans la cavité.

5. Champ et potentiel créé par un noyau non uniformément chargé, d'après ENAC 2005

Les noyaux de certains atomes légers peuvent être modélisés par une distribution volumique de charge à l'intérieur d'une sphère de centre O et de rayon a . La densité volumique de charges est donnée par $\rho = \rho_0 \left(1 - \frac{r^2}{a^2}\right)$ pour $r < a$ avec ρ_0 une constante positive.

1. Calculer la charge totale Q du noyau.
2. Étudier les invariances et les symétries de la distribution. En déduire l'orientation et les dépendances du champ électrique créé par cette distribution.
3. Déterminer l'expression du champ électrique à l'extérieur du noyau.
4. Même question à l'intérieur du noyau.
5. En prenant l'origine des potentiels à l'infini, calculer le potentiel électrique à l'extérieur du noyau.
6. Même question à l'intérieur.

6. Sphère chargée non uniformément

Soit une boule de rayon R , chargée en volume.

Déterminer la distribution volumique de charges $\rho(r)$ pour que le champ électrostatique soit de module constant E_0 à l'intérieur de la boule.

Calculer le champ électrostatique et le potentiel électrostatique en tout point de l'espace.

7. Ruban infini

On considère un ruban infini selon l'axe Ox , de largeur b selon Oy (symétrique par rapport à cet axe), et d'épaisseur nulle. Il est chargé uniformément en surface avec une densité σ .

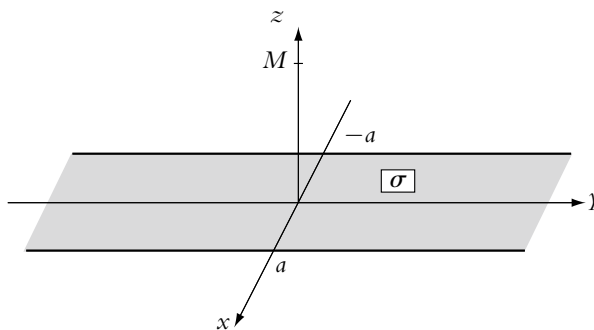


Figure 45.13

On veut calculer le champ en tout point M de l'axe Oz .

1. Déterminer la direction du champ en M par l'étude des symétries de la distribution des charges.

2. On décompose le ruban en fil infini de largeur dy et de densité linéique λ .
 - a) Déterminer la relation entre λ , σ et dy .
 - b) Déterminer le champ élémentaire $d\vec{E}$ créé par l'un de ces fils en M .
 - c) Calculer le champ créé en M par le ruban.

8. **Lame chargée**

Des charges électriques positives sont distribués uniformément dans le volume compris entre deux plan infinis orthogonaux à un axe Oz de l'espace, de cotes respectives $z = +a$ et $z = -a$.

1. A l'aide du théorème de Gauss, calculer le champ $\vec{E}(z)$ en tout point de l'espace.
2. Calculer le potentiel $V(z)$ en tout point de l'espace.

9. **Deux lames de charges opposées**

Soit la distribution volumique de charges définie en coordonnées cartésiennes par :

$$\rho = \begin{cases} +\rho_0 & \text{si } 0 \leq x \leq a \\ -\rho_0 & \text{si } -a \leq x \leq 0 \end{cases}$$

1. Étudier les symétries et les invariances de la distribution.
2. Calculer le champ électrostatique créé en tout point de l'espace par cette distribution.
3. En déduire l'expression du potentiel électrostatique créé en tout point par cette distribution. On notera V_0 la valeur du potentiel en $x = 0$.

B. Exercices et problèmes

1. **Modélisation d'un nuage mince**

On considère, dans le vide et loin de toute charge, un nuage électrique plan Π suffisamment étendu pour le considérer d'épaisseur négligeable, d'équation $z = h$ et chargé avec la densité surfacique constante négative σ . On donne : $\frac{1}{4\pi\epsilon_0} = 9.10^9$ U.S.I.

1. Déterminer vectoriellement le champ électrique en tout point M de l'espace, en fonction de z , en utilisant le théorème de Gauss.
2. On considère maintenant que le plan xOy représentant le sol porte une densité de charge opposée à celle du nuage. Déterminer le champ total (sol + nuage) pour $0 < z < h$.
3. Déterminer le potentiel total (sol + nuage) pour $0 < z < h$ sachant que $V(0) = 0$.
4. Ce condensateur modélise un nuage carré de 10 km de côté à une hauteur $h = 2$ km. Déterminer la capacité du condensateur ainsi formé (application numérique).
5. On considère maintenant un nuage d'orage. Sachant que lorsque l'éclair se forme, le champ a pour valeur dans ce cas 25 kV.m^{-1} , déterminer le potentiel V_0 du nuage.

2. Deux conducteurs cylindriques parallèles, d'après ENSAIT PC 2003

Soit un conducteur rectiligne de rayon $a = 5,0 \text{ mm}$ et de longueur infinie portant une charge électrique positive de densité volumique $\rho = 8,0 \cdot 10^{-9} \text{ C.m}^{-3}$.

1. Calculer le champ électrique créé par ce conducteur en tout point de l'espace.
2. En déduire le potentiel électrique en tout point de l'espace en prenant l'origine des potentiels pour $r = a$. On justifiera qu'on ne prend pas l'origine à l'infini.
3. En déduire la différence de potentiel ΔU entre la surface du conducteur et un point extérieur au conducteur à une distance $x = 1,0 \text{ cm}$ de l'axe du conducteur.
4. On place deux conducteurs de ce type l'un chargé positivement, l'autre négativement avec des densités volumiques de charge opposées à une distance D l'un de l'autre et parallèles entre eux. Exprimer le potentiel en un point M situé entre les deux conducteurs à une distance x du conducteur chargé négativement. On prendra comme origine des potentiels le plan équidistant des deux conducteurs.
5. Tracer la courbe donnant ce potentiel en fonction de x .

3. Analogie électrostatique - gravitation

1. Soit une demi-sphère de centre O et de rayon R uniformément chargée avec une densité volumique de charges ρ . Déterminer les invariances du système et en déduire les conséquences sur le champ électrostatique qu'il crée.
2. Mêmes questions pour les plans de symétrie et d'antisymétrie.
3. En déduire l'orientation du champ électrostatique créé par cette distribution en O .
4. Déterminer l'expression du champ électrostatique créé par cette distribution en O à l'aide d'un calcul direct.
5. On considère maintenant un disque uniformément chargé.
Déterminer l'expression du champ électrostatique créé par ce disque en un point de son axe.
6. En utilisant ce résultat, retrouver l'expression du champ électrostatique créé par une demi-sphère de centre O et de rayon R uniformément chargée avec une densité volumique de charges ρ en son centre O .
7. Expliciter les analogies entre électrostatique et gravitation.
8. On modélise un lac par une demi-sphère de centre O et de rayon R rempli d'eau de masse volumique ρ_1 uniforme. On suppose que le sol a pour masse volumique ρ_0 et que la masse y est uniformément répartie. On note $\vec{g}_0$ le champ de gravitation de la Terre au niveau du sol loin du lac et $\vec{g}_1$ celui au centre O du lac.
Calculer $\Delta g = g_1 - g_0$.
9. Application numérique : $\rho_1 = 3,0 \cdot 10^3 \text{ kg.m}^{-3}$, $\rho_0 = 1,0 \cdot 10^3 \text{ kg.m}^{-3}$ et $R = 1,0 \text{ km}$.
Donner l'ordre de grandeur de Δg .

4. Champ créé par un cône ; champ gravitationnel créé par un volcan

On désire calculer le champ et le potentiel au sommet du cône.

1. a) Déterminer le champ électrostatique créé par un disque de centre O , d'épaisseur nulle, de rayon r , de densité surfacique de charge σ uniforme en tout point de son axe de symétrie Oz .
- b) Même question avec le potentiel électrostatique
- c) On considère maintenant que le disque a une épaisseur dz et une densité volumique de charge uniforme ρ . À partir des résultats précédents, déterminer le champ élémentaire dE et le potentiel élémentaire dV en tout point de l'axe Oz .
2. On considère un cône renversé de hauteur h , de base de rayon R , de demi-angle au sommet α . Il possède une densité volumique de charge constante ρ .
 - a) On considère ce cône comme un empilement de disques d'épaisseur dz . Déterminer le champ en S sommet du cône.
 - b) Déterminer le potentiel au sommet S du cône
3. a) Exprimer par analogie entre la force de Coulomb et celle de Newton, le champ gravitationnel créé par une masse ponctuelle.
- b) On assimile la Terre à une boule homogène de masse $M_T = 6.10^{24}$ kg, de rayon $R_T = 6400$ km et de masse volumique moyenne uniforme $\mu = 5.10^3$ kg.m⁻³.
Déterminer le champ gravitationnel créé par la Terre à une distance $r > R_T$ de son centre.
- c) On considère un volcan assimilable à un cône renversé de hauteur $h = 3200$ m, de rayon de base $R_b = 20$ km (Etna). Déterminer le champ de gravitation total (terre +volcan), au sommet de celui-ci.
- d) Calculer la valeur numérique totale du champ. Comparer à la valeur du champ en l'absence du volcan.

5. Champ et potentiel d'une jonction, d'après ENSTIM 1995

Le germanium pur (Ge) est un bon isolant électrique. Lorsqu'on introduit des impuretés en très faible concentration, par exemple de l'antimoine (Sb), la conductivité électrique du germanium augmente fortement : on obtient un "semi-conducteur dopé", noté Ge : Sb, dont les propriétés électriques dépendent à la fois du nombre d'atomes Sb introduit par unité de volume, noté N et de la température absolue T .

Lorsqu'un semi-conducteur présente, dans une région très localisée de l'espace, une variation très brutale de la concentration en dopant, voire un changement de la nature du dopant, on dit qu'on a une jonction. Au voisinage de la jonction, dans une région dite "zone de charge", le cristal acquiert une distribution de charge électrique non nulle que l'on se propose d'étudier. Les propriétés qui en résultent sont à la base de la caractéristique des diodes, des transistors, et de tous les circuits intégrés.

1. On considère un plan infini d'équation $z = 0$, portant une densité surfacique de charge σ constante. ce plan est plongé dans un milieu quasi isolant dans lequel la permittivité électrique du vide ϵ_0 doit être simplement remplacée par le produit $\epsilon_0 \epsilon_r$, où ϵ_r est la permittivité relative du milieu considéré.
- a) Étudier et justifier soigneusement les propriétés de symétrie du champ électrique créé par la distribution de charge.
- b) Déterminer le champ électrique en un point M quelconque.

2. On se place dans le germanium, de permittivité relative ϵ_r , et on suppose que la densité volumique de charge ρ , autour d'une jonction située dans le plan $z = 0$, a l'allure de la figure (45.14).

- si $z < -L_1$: $\rho = 0$
- si $-L_1 < z < 0$: $\rho = \rho_1 < 0$
- si $0 < z < L_2$: $\rho = \rho_2 > 0$
- si $z > L_2$: $\rho = 0$

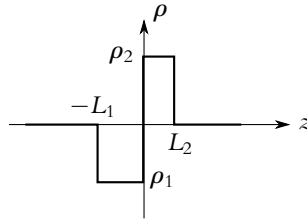


Figure 45.14

La jonction est suffisamment large pour supposer que la distribution de charge est totalement invariante par toute translation dans le plan Oxy .

a) Sachant que la distribution de charge est globalement neutre, établir la relation vérifiée par L_1 , L_2 , ρ_1 et ρ_2 .

b) Donner l'expression de l'élément de champ dE créé en tout point M de l'espace, par une couche élémentaire d'épaisseur dz , à la cote h ($-L_1 < h < L_2$). On utilisera les résultats de la question (1).

c) Déterminer le champ électrique en tout point M de l'espace créé par la distribution précédente. On distinguera les quatre régions de l'espace suivant la valeur de z .

Représenter graphiquement E en fonction de z .

d) En déduire le potentiel $V(M)$. On choisira l'origine des potentiels dans le plan $z = 0$.

Représenter graphiquement V en fonction de z .

e) Donner l'expression de la différence de potentiel V_0 entre deux points situés de part et d'autre de la zone de charge.

3. La région $z > 0$ a été dopée avec de l'antimoine à raison de $N_2 = 1,6 \cdot 10^{21}$ atomes Sb par m^3 , tandis que la région $z < 0$ a été dopée avec du bore avec un nombre d'atomes par unité de volume $N_1 \gg N_2$.

On admet que, dans la zone de charge, chaque atome Sb est ionisé en Sb^+ . Les électrons ainsi libérés traversent spontanément la plan $z = 0$, et chaque atome B situé dans la zone de charge capte un électron, se transformant en ion B^- .

a) En déduire ρ_1 et ρ_2 en fonction de N_1 et N_2 .

b) Le système ainsi constitué est une diode à jonction dont la tension de seuil est voisine de V_0 . En déduire une expression approchée de la largeur δ de la zone de charge.

c) Application numérique : calculer δ si $V_0 = 0,3 \text{ V}$; $\epsilon_0 = 8,84 \cdot 10^{-12} \text{ USI}$; $\epsilon_r = 16$ et $e = 1,6 \cdot 10^{-19} \text{ C}$.

6. Anneau de stockage pour molécules polaires, d'après Polytechnique
PC 2005

On étudie la possibilité de guider le mouvement de molécules polaires avec un système électrostatique formé de six électrodes cylindriques et parallèles $C_i, i = 1, 2, \dots, 6$ disposées aux sommets d'un hexagone régulier auquel elles sont orthogonales (figure 45.15).

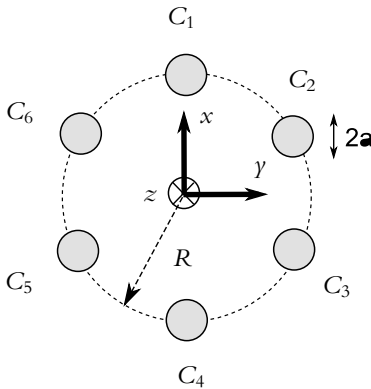


Figure 45.15

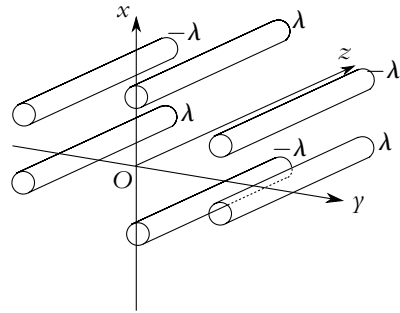


Figure 45.16

Leur rayon a est très inférieur au côté R de l'hexagone, $a \ll R$. Elles portent des densités linéiques de charge égales alternativement à λ ($\lambda > 0$) pour les électrodes impaires et $-\lambda$ pour les paires ; on considèrera que ces charges sont fixes et uniformément réparties à leur surface. On négligera les effets d'extrémités, l'ensemble pouvant être considéré comme invariant par translation selon l'axe central Oz du système. On utilisera un système de coordonnées cylindriques (r, θ, z) avec comme repère orthonormé $(\vec{u}_r, \vec{u}_\theta, \vec{u}_z)$.

1. Analyse des symétries

- Quelles conclusions sur le champ $\vec{E}$ et le potentiel électrostatique V tire-t-on de l'invariance par translation du système ?
- Considérer la symétrie par rapport à un plan perpendiculaire à l'axe. Quelle propriété du champ électrique $\vec{E}$ en déduit-on ?
- Même question pour l'un des trois plans passant par l'axe central et les axes de deux électrodes opposées.
- Montrer que les trois plans passant par l'axe et à égale distance des électrodes sont équipotentiels.
- Quelle est la période angulaire d'invariance du système par rotation autour de l'axe Oz ? En déduire une expression générale du potentiel $V(r, \theta, z)$ sous forme d'une série.

2. Soit une électrode de densité linéique de charge λ . Déterminer le champ électrostatique créé par cette électrode en un point P à l'aide de la distance D de ce point à son axe ($D > a$). En déduire une expression du potentiel électrostatique correspondant.

3. On considère maintenant l'ensemble des électrodes du système. Montrer que, en le choisissant nul sur l'axe central, le potentiel électrostatique en un point P est donné par l'expression :

$$V(P) = \frac{\lambda}{2\pi\epsilon_0} \ln \left(\frac{D_2 D_4 D_6}{D_1 D_3 D_5} \right)$$

où D_i désigne la distance de P à l'axe de l'électrode C_i .

4. Pour expliciter le potentiel en fonction des coordonnées de P , il est commode de considérer le plan xOy comme plan de représentation des nombres complexes. Le point P y est repéré par $Z = x + iy = r \exp(i\theta)$, les axes des électrodes impaires le sont par $(R, jR, j^2 R)$ et ceux des électrodes paires par $(-R, -jR, -j^2 R)$, avec $j = \exp(i2\pi/3)$ racine cubique de l'unité. Montrer que :

$$\frac{D_2 D_4 D_6}{D_1 D_3 D_5} = \left| \frac{R^3 + \overline{Z}^3}{R^3 - \overline{Z}^3} \right|$$

5. On s'intéresse à la partie centrale $r \ll R$. Montrer que le potentiel électrostatique y est donné par :

$$V(r, \theta, z) = \frac{\lambda}{\pi\epsilon_0} \left(\frac{r}{R} \right)^3 \cos(3\theta)$$

Cette expression respecte-t-elle les symétries étudiées en question (1) ?

6. Déterminer les potentiels V_0 des électrodes impaires dans l'hypothèse $a \ll R$ en fonction de R , a et λ . Quel est celui des électrodes paires ?

7. On considère le système comme un condensateur, les trois électrodes impaires formant l'une des armatures, les trois paires l'autre. Déterminer la capacité par unité de longueur correspondante C .

Montrer que le potentiel électrostatique dans la partie centrale de l'hexapôle s'exprime simplement en fonction de cette capacité linéique et de la tension V_0 .

8. Application numérique : calculer la capacité électrostatique par unité de longueur d'un hexapôle ayant $R = 2,5$ cm et $a = 2,5$ mm.

Dipôle électrostatique (MPSI, PCSI)

46

Pour toutes les distributions envisagées jusqu'à présent, on n'a fait aucune approximation et on a calculé le champ et le potentiel électrostatiques. Dans ce chapitre, on s'intéresse au potentiel et au champ électrostatiques « loin » de la distribution, à savoir à très grande distance de la distribution par rapport aux dimensions caractéristiques de cette dernière. Pour cela, on va définir la notion de dipôle électrostatique et obtenir les expressions dans le cadre de l'approximation énoncée.

1. Définition, potentiel et champ créés

1.1 Définition du dipôle électrostatique

On appelle *dipôle électrostatique* le système constitué de deux charges ponctuelles opposées $-q$ et q situées en deux points N et P distants de a et tels que $a = NP$ soit très petite devant les autres distances envisagées.

1.2 Moment dipolaire

On définit alors le *moment dipolaire* de la distribution par le vecteur :

$$\vec{p} = q\overrightarrow{NP}$$

L'unité de cette quantité est donc, dans les unités du système international, le coulomb mètre (C.m).

Compte tenu des ordres de grandeur (Cf. paragraphe 3.6), on préfère souvent donner les moments dipolaires en debye de symbole D : $1 \text{ D} = 3,336.10^{-30} \text{ C.m}$.

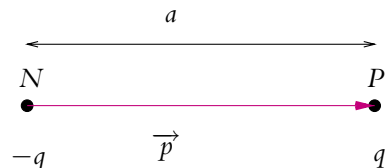


Figure 46.1 Moment dipolaire.

On peut aussi noter qu'un dipôle électrostatique est la limite d'un ensemble de deux charges ponctuelles $-q$ et q placées en N et P tel que la distance NP tend vers zéro tandis que le moment dipolaire $p = qNP$ reste constant et fini.

1.3 Potentiel électrostatique créé par un dipôle

a) Invariances et symétries du dipôle

On utilise les coordonnées sphériques pour décrire le problème. La distribution est invariante par rotation autour de l'axe du dipôle, donc le champ électrostatique $\vec{E}(M)$ et le potentiel électrostatique $V(M)$ ne dépendent pas de φ .

Pour tenir compte de ces propriétés, on se place en coordonnées sphériques dans un plan méridien ou plan tel que $\varphi = \text{constante}$ et on utilise les coordonnées polaires dans ce plan (c'est-à-dire les deux autres coordonnées sphériques r et θ).

b) Calcul du potentiel créé par un dipôle

Le potentiel créé en M par le dipôle est égal d'après ce qu'on a écrit précédemment à :

$$V(M) = \frac{1}{4\pi\epsilon_0} \frac{-q}{MN} + \frac{1}{4\pi\epsilon_0} \frac{q}{MP} = \frac{q}{4\pi\epsilon_0} \left(\frac{1}{MP} - \frac{1}{MN} \right)$$

On appelle O le milieu du segment $[NP]$.

On se place en coordonnées polaires d'origine O dans un plan contenant le point M et la droite (NP) sans perdre en généralité du fait de la symétrie autour de cette droite (Cf. figure 46.2 pour les notations).

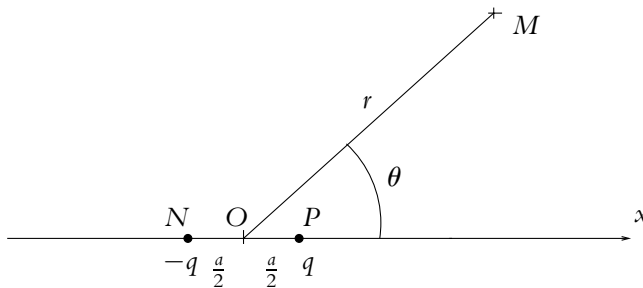


Figure 46.2 Moment dipolaire.

On obtient :

$$(MP)^2 = (\vec{MO} + \vec{OP})^2 = MO^2 + OP^2 + 2\vec{MO} \cdot \vec{OP} = r^2 + \frac{a^2}{4} - 2r\frac{a}{2} \cos \theta$$

D'où

$$\frac{1}{MP} = \frac{1}{r} \left(1 - \frac{a}{r} \cos \theta + \frac{a^2}{4r^2} \right)^{-\frac{1}{2}}$$

Comme le point M est « loin » de la distribution, on a $a \ll r$ et on peut effectuer un développement limité de la relation précédente :

$$\frac{1}{MP} = \frac{1}{r} \left(1 + \frac{1}{2} \frac{a}{r} \cos \theta + o\left(\frac{a}{r}\right) \right)$$

De la même manière, on obtient :

$$\frac{1}{MN} = \frac{1}{r} \left(1 - \frac{1}{2} \frac{a}{r} \cos \theta + o\left(\frac{a}{r}\right) \right)$$

Ce qui permet d'obtenir l'expression du potentiel dans le cadre de l'approximation considérée (au premier ordre en $\frac{a}{r}$) :

$$V(M) = \frac{q}{4\pi\epsilon_0} \frac{1}{r} \left(1 + \frac{1}{2} \frac{a}{r} \cos \theta - 1 + \frac{1}{2} \frac{a}{r} \cos \theta \right) = \frac{qa \cos \theta}{4\pi\epsilon_0 r^2}$$

Or $qa = p$ et $p \cos \theta = \vec{p} \cdot \frac{\vec{OM}}{OM}$ donc

$$V(M) = \frac{\vec{p} \cdot \vec{OM}}{4\pi\epsilon_0 OM^3} = \frac{p \cos \theta}{4\pi\epsilon_0 r^2}$$

1.4 Champ électrostatique créé par un dipôle

a) Invariances et symétries du dipôle

Du fait des invariances déjà étudiées et de l'indépendance par rapport à φ , on se place dans un plan contenant le point M et l'axe du dipôle et on utilise les coordonnées polaires dans ce plan.

De plus, tout plan contenant cet axe est un plan de symétrie des sources donc le champ électrostatique $\vec{E}(M)$ appartient à ce plan.

b) Expression du champ créé

On déduit l'expression du champ électrostatique du potentiel à partir de la relation :

$$\vec{E} = -\vec{\text{grad}} V$$

et de l'expression obtenue pour le potentiel :

$$V(M) = \frac{p \cos \theta}{4\pi\epsilon_0 r^2}$$

En coordonnées polaires :

$$\begin{cases} E_r = -\frac{\partial V}{\partial r} = \frac{2p \cos \theta}{4\pi\epsilon_0 r^3} \\ E_\theta = -\frac{1}{r} \frac{\partial V}{\partial \theta} = \frac{p \sin \theta}{4\pi\epsilon_0 r^3} \end{cases}$$

Donc

$$\vec{E}(M) = \frac{1}{4\pi\epsilon_0} \frac{2p \cos \theta}{r^3} \vec{u}_r + \frac{1}{4\pi\epsilon_0} \frac{p \sin \theta}{r^3} \vec{u}_\theta$$

Ce champ peut également s'écrire :

$$\vec{E}(M) = \frac{1}{4\pi\epsilon_0} \frac{3 \left(\vec{p} \cdot \overrightarrow{OM} \right) \overrightarrow{OM} - OM^2 \vec{p}}{OM^5}$$

Il suffit par exemple d'écrire que $3 = 2 - 1$ et d'identifier $p \cos \theta$ au produit scalaire de $\vec{p}$ par le vecteur unitaire directeur de $\overrightarrow{OM}$ et $p \cos \theta \vec{u}_r - p \sin \theta \vec{u}_\theta$ à la projection de $\vec{p}$ dans la base $(\vec{u}_r, \vec{u}_\theta)$.

Le potentiel créé par le dipôle décroît en $\frac{1}{r^2}$ et le champ en $\frac{1}{r^3}$ alors que pour une charge ponctuelle, ils décroissent respectivement en $\frac{1}{r}$ et $\frac{1}{r^2}$: les effets d'un dipôle se font ressentir à moins grande distance que ceux d'une charge seule.

Champ et potentiel électrostatique s'expriment en fonction du moment dipolaire $\vec{p}$ et non en fonction de a et de q séparément : le moment dipolaire est la grandeur caractéristique du dipôle.

1.5 Lignes de champ et équipotentielles

a) Lignes de champ

On cherche les lignes tangentes au champ électrostatique en tout point en utilisant la relation :

$$\vec{E}(M) \wedge d\overrightarrow{OM} = \vec{0}$$

qu'on écrit en coordonnées polaires :

$$\begin{pmatrix} E_r \\ E_\theta \\ 0 \end{pmatrix} \wedge \begin{pmatrix} dr \\ r d\theta \\ 0 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix}$$

soit :

$$rE_r d\theta = E_\theta dr$$

On obtient :

$$\frac{2p \cos \theta}{4\pi\epsilon_0 r^2} d\theta = \frac{p \sin \theta}{4\pi\epsilon_0 r^3} dr$$

En séparant les variables, on en déduit :

$$\frac{dr}{r} = 2 \frac{\cos \theta d\theta}{\sin \theta} = 2 \frac{d(\sin \theta)}{\sin \theta}$$

Soit après intégration :

$$\ln r = 2 \ln |\sin \theta| + C$$

ou en prenant l'exponentielle :

$$r = \lambda \sin^2 \theta$$

où λ est une constante.

D'où l'allure des lignes de champ :

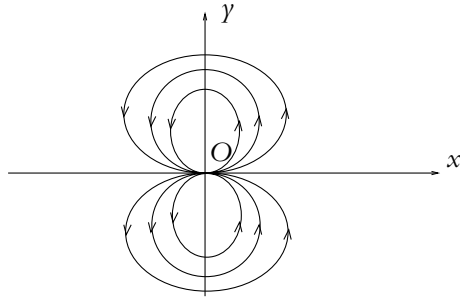


Figure 46.3 Lignes de champ du dipôle électrostatique.

b) Équipotentielles

Elles sont définies par l'ensemble des points ayant le même potentiel soit

$$V = V_0 = \frac{p \cos \theta}{4\pi\epsilon_0 r^2} \quad \text{donc} \quad r = \mu \sqrt{|\cos \theta|}$$

où μ est une constante.

D'où l'allure des équipotentielle :

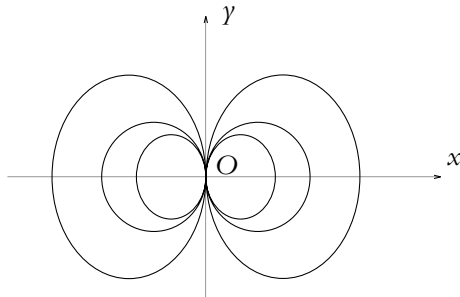


Figure 46.4 Équipotentielle du dipôle électrostatique.

c) Tracé du diagramme électrique

On peut superposer les deux tracés précédents sur un même diagramme électrique.

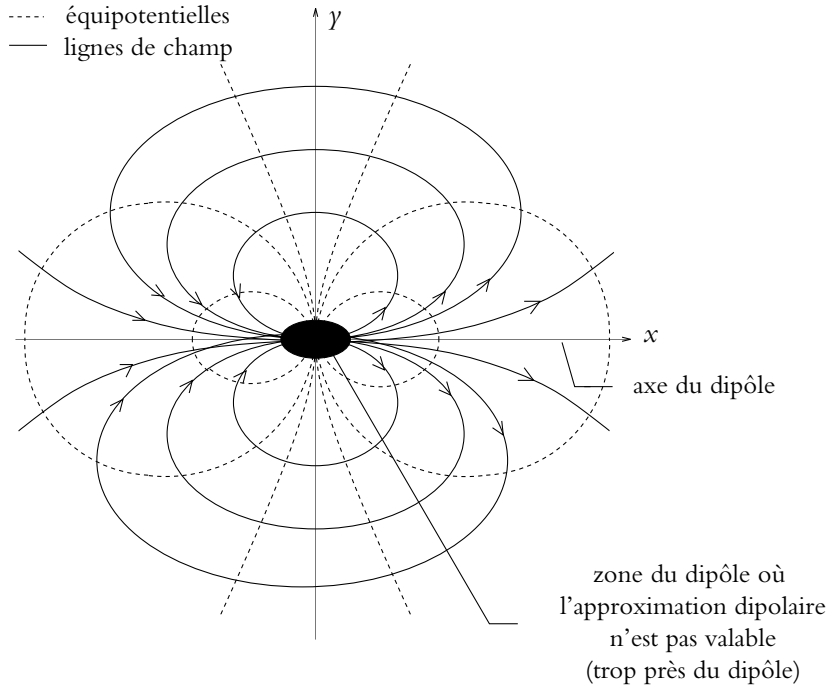


Figure 46.5 Diagramme électrique du dipôle électrostatique.

► **Remarques :** On peut formuler les remarques suivantes :

- Lignes de champ et équipotentiels sont en tout point de l'espace orthogonales.
- La zone où se trouve le dipôle est une zone où les calculs précédents ne sont pas valables, ils ne donnent pas l'allure des lignes de champ ou des équipotentiels : on n'est pas suffisamment « loin » de la distribution. Cette zone a été noircie sur la représentation.

2. Action d'un champ extérieur sur un dipôle

2.1 Cas d'un champ uniforme

On s'intéresse tout d'abord à l'action subie par un dipôle dans un champ électrostatique extérieur uniforme en déterminant la force et le moment qu'il subit.

a) Force exercée sur un dipôle par un champ électrostatique extérieur uniforme

La force qui s'exerce sur un dipôle électrostatique est la somme des forces qui s'exercent sur chacune des deux charges. On peut noter que la force d'interaction entre les deux charges constituant le dipôle donne une contribution nulle par le principe des actions réciproques.

Donc

$$\vec{F} = q\vec{E}(P) - q\vec{E}(N)$$

Si le champ est uniforme, $\vec{E}(P) = \vec{E}(N)$ et $\vec{F} = \vec{0}$.

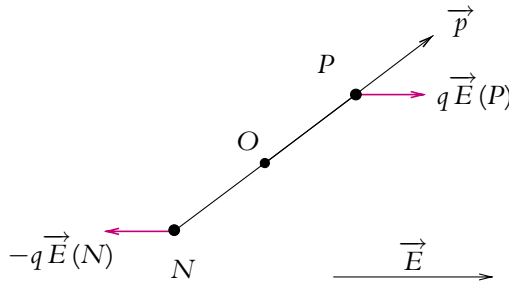


Figure 46.6 Action d'un champ électrostatique uniforme sur un dipôle.

Un champ électrostatique uniforme n'exerce pas de force sur un dipôle électrostatique.

b) Moment exercé sur un dipôle par un champ électrostatique extérieur uniforme

De la même manière, le moment qui s'exerce sur le dipôle est :

$$\vec{M}_O = \vec{OP} \wedge q\vec{E}(P) - \vec{ON} \wedge q\vec{E}(N)$$

$$\vec{M}_O = q(\vec{OP} - \vec{ON}) \wedge \vec{E}$$

$$\vec{M}_O = q\vec{NP} \wedge \vec{E}$$

$$\vec{M}_O = \vec{p} \wedge \vec{E}$$

On notera que le moment obtenu est indépendant du point où on le calcule. On le notera $\vec{\Gamma}$

c) Analyse qualitative de l'action d'un champ électrostatique extérieur uniforme sur un dipôle

L'action d'un champ électrostatique extérieur uniforme sur un dipôle se réduit à un couple de moment :

$$\vec{\Gamma} = \vec{p} \wedge \vec{E}$$

En notant θ l'angle entre $\vec{p}$ et $\vec{E}$ et $\vec{k}$ un vecteur unitaire perpendiculaire à $\vec{p}$ et $\vec{E}$, le couple s'écrit :

$$\vec{\Gamma} = pE \sin \theta \vec{k}$$

D'après le théorème du moment cinétique, on aura un état d'équilibre si le moment des forces est nul donc si $\sin \theta = 0$, à savoir si

$$\theta = 0 \quad \text{ou} \quad \theta = \pi$$

Cela correspond à des positions du dipôle parallèle ou antiparallèle au champ électrostatique.

Pour discuter de la stabilité de ces positions d'équilibre, on écarte légèrement le dipôle de sa position d'équilibre et on analyse le type de mouvement que le moment du couple tend à créer.

- Dipôle parallèle :

Les forces tendent à ramener le dipôle dans sa position d'équilibre donc l'équilibre des dipôles parallèles au champ est stable.

- Dipôle antiparallèle :

Les forces tendent à écarter le dipôle de sa position d'équilibre donc l'équilibre des dipôles antiparallèles au champ est instable.

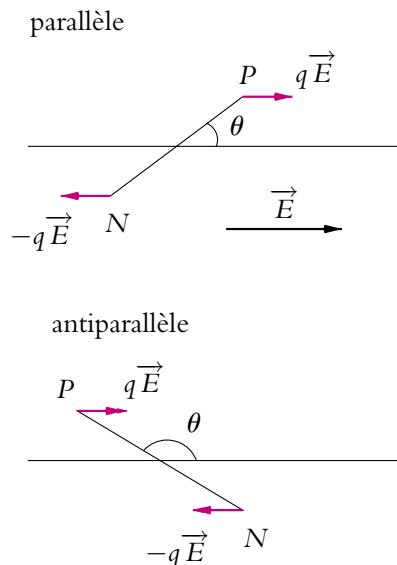


Figure 46.7 Etude de la stabilité des positions d'équilibre d'un dipôle dans un champ électrostatique extérieur.

Un champ électrostatique uniforme tend donc à orienter les dipôles électrostatiques suivant les lignes de champ.



À l'échelle du dipôle, tout champ est en première approximation uniforme : l'effet principal d'un champ extérieur sur un dipôle est son orientation suivant les lignes de champ.

2.2 Cas d'un champ non uniforme

Aucune expression établie dans ce paragraphe n'est exigible dans le cadre du programme ; seule l'analyse qualitative du phénomène est à retenir. En cas de besoin, il faudra démontrer toutes les relations qui pourront être obtenues ici.

a) Force exercée sur un dipôle par un champ électrostatique extérieur non uniforme

La force s'exprime par :

$$\vec{F} = q \vec{E}(P) - q \vec{E}(N) = q \left(\vec{E}(P) - \vec{E}(N) \right)$$

Le champ extérieur n'étant pas uniforme, cette force n'est pas nulle : il faut tenir compte des variations du champ entre P et N . Or la définition d'un dipôle impose à la distance PN d'être faible devant les distances caractéristiques du problème donc en particulier devant la distance caractéristique des variations du champ électrostatique extérieur auquel le dipôle est soumis. Par conséquent, on peut effectuer un développement limité du champ électrostatique en P et en N .

On va se placer en coordonnées cartésiennes pour effectuer la démonstration. Le résultat est cependant tout à fait général et indépendant du système de coordonnées dans lequel on se place. On cherchera donc à obtenir une expression intrinsèque. Soient (x, y, z) les coordonnées du milieu de $[PN]$ et $(\Delta x, \Delta y, \Delta z)$ les composantes de $\vec{NP}$. Les coordonnées des points P et N sont donc respectivement : $\left(x + \frac{\Delta x}{2}, y + \frac{\Delta y}{2}, z + \frac{\Delta z}{2}\right)$ et $\left(x - \frac{\Delta x}{2}, y - \frac{\Delta y}{2}, z - \frac{\Delta z}{2}\right)$.

On en déduit pour la composante suivant Ox de la force :

$$F_x = q \left(E_x \left(x + \frac{\Delta x}{2}, y + \frac{\Delta y}{2}, z + \frac{\Delta z}{2} \right) - E_x \left(x - \frac{\Delta x}{2}, y - \frac{\Delta y}{2}, z - \frac{\Delta z}{2} \right) \right)$$

Le développement limité au premier ordre de la composante sur Ox du champ donne :

$$\begin{aligned} F_x &= q \left(E_x(x, y, z) + \frac{\Delta x}{2} \frac{\partial E_x}{\partial x} + \frac{\Delta y}{2} \frac{\partial E_x}{\partial y} + \frac{\Delta z}{2} \frac{\partial E_x}{\partial z} + \right. \\ &\quad \left. - E_x(x, y, z) + \frac{\Delta x}{2} \frac{\partial E_x}{\partial x} + \frac{\Delta y}{2} \frac{\partial E_x}{\partial y} + \frac{\Delta z}{2} \frac{\partial E_x}{\partial z} \right) \\ &= q \left(\Delta x \frac{\partial E_x}{\partial x} + \Delta y \frac{\partial E_x}{\partial y} + \Delta z \frac{\partial E_x}{\partial z} \right) \\ &= \vec{p} \cdot \left(\vec{\text{grad}} E_x \right) \end{aligned}$$

Par un raisonnement analogue, on obtient :

$$F_y = \vec{p} \cdot \left(\vec{\text{grad}} E_y \right) \quad \text{et} \quad F_z = \vec{p} \cdot \left(\vec{\text{grad}} E_z \right)$$

On note $\vec{F} = \left(\vec{p} \cdot \overrightarrow{\text{grad}} \right) \vec{E}$, ce qui signifie que $F_x = \vec{p} \cdot \left(\overrightarrow{\text{grad}} E_x \right)$, $F_y = \vec{p} \cdot \left(\overrightarrow{\text{grad}} E_y \right)$ et $F_z = \vec{p} \cdot \left(\overrightarrow{\text{grad}} E_z \right)$ en **coordonnées cartésiennes**.

L'expression de l'opérateur $\vec{p} \cdot \overrightarrow{\text{grad}}$ est beaucoup plus compliquée dans les autres systèmes de coordonnées, elle sera donnée si besoin est.

Cas d'un dipôle rigide

Un dipôle est dit *rigide* si la distance entre les charges positive et négative constituant le dipôle est constante en toutes circonstances.

Dans ce cas, on peut « rentrer » les termes $q\Delta x$, $q\Delta y$ et $q\Delta z$ dans la dérivée partielle. On obtient ainsi l'expression :

$$\vec{F} = \overrightarrow{\text{grad}} \left(\vec{p} \cdot \vec{E} \right)$$

Cas d'un dipôle non rigide

Pour un dipôle rigide en champ non uniforme, on a obtenu pour la composante F_x de la force :

$$F_x = q \left(\Delta x \frac{\partial E_x}{\partial x} + \Delta y \frac{\partial E_x}{\partial y} + \Delta z \frac{\partial E_x}{\partial z} \right)$$

Le fait que le dipôle ne soit plus rigide implique un mouvement relatif des points N et P . Cela aura pour conséquence de modifier les valeurs de Δx , Δy et Δz , mais ne changera pas le résultat final :

$$\vec{F} = \left(\overrightarrow{p(E)} \cdot \overrightarrow{\text{grad}} \right) \vec{E}$$

On doit simplement tenir compte du fait que le dipôle dépend du champ électrique.

b) Moment exercé sur un dipôle par un champ électrostatique extérieur non uniforme

On pose $\vec{E}(P) = \vec{E}(O) + \delta \vec{E}(P)$ et $\vec{E}(N) = \vec{E}(O) + \delta \vec{E}(N)$.

On procède pour le moment comme pour la force en limitant les développements limités au premier ordre non nul :

$$\vec{M}_O = \overrightarrow{OP} \wedge q \vec{E}(P) + \overrightarrow{ON} \wedge (-q \vec{E}(N))$$

$$\vec{M}_O = q \overrightarrow{OP} \wedge \left(\vec{E}(O) + \delta \vec{E}(P) \right) - q \overrightarrow{ON} \wedge \left(\vec{E}(O) + \delta \vec{E}(N) \right)$$

$$\vec{M}_O = q \left(\overrightarrow{OP} - \overrightarrow{ON} \right) \wedge \vec{E}(O) + q \left(\overrightarrow{OP} \wedge \delta \vec{E}(P) - \overrightarrow{ON} \wedge \delta \vec{E}(N) \right)$$

$$\vec{M}_O \simeq q \overrightarrow{NP} \wedge \vec{E}(O)$$

$$\vec{M}_O = \vec{p} \wedge \vec{E}(O)$$

en négligeant les termes du premier ordre par la suite pour ne garder que le terme d'ordre zéro qui est non nul.

Donc :

$$\vec{\mathcal{M}}_o = \vec{p} \wedge \vec{E}$$

c) Analyse qualitative de l'action d'un champ électrostatique extérieur non uniforme sur un dipôle rigide

L'action d'un champ extérieur non uniforme dépend à la fois de la force et du couple qui s'exercent sur le dipôle.

L'action du couple tend (comme dans le cas du champ uniforme) à orienter le dipôle dans le sens des lignes de champ. Ce sera l'action principale d'un champ extérieur, qu'il soit ou non uniforme.

Quand le champ n'est pas uniforme, on doit ajouter l'action de la force :

$$\vec{F} = (\vec{p} \cdot \vec{\text{grad}}) \vec{E}$$

Dans le cas particulier où le dipôle est parallèle au champ, on doit distinguer :

- le cas où le champ et le dipôle sont orientés dans le même sens : la force tend à attirer le dipôle vers les champs intenses,
- le cas où le champ et le dipôle sont orientés dans des sens opposés : la force tend à attirer le dipôle vers les champs faibles.



Il n'y a aucune raison *a priori* pour que dipôle et champ soient colinéaires, si ce n'est que le champ tend à orienter le dipôle suivant les lignes de champ (à cause du couple). Les dipôles seront alors attirés par les champs intenses. D'autre part, il ne faut pas oublier le phénomène d'agitation thermique qui fait osciller le dipôle autour de la position de colinéarité.

2.3 Énergie potentielle d'un dipôle dans un champ électrostatique extérieur (PCSI)

On s'intéresse ici à l'énergie électrostatique d'un dipôle dans un champ extérieur et non à l'énergie propre du dipôle.

a) Expression de l'énergie potentielle d'un dipôle rigide dans un champ électrostatique extérieur

On suppose en outre dans ce paragraphe et le suivant que le dipôle est rigide.

Soit V le potentiel associé au champ dans lequel on place le dipôle.

L'énergie potentielle d'interaction du dipôle avec le champ extérieur s'écrit donc :

$$E_p = \sum_i q_i V(P_i) = qV(P) - qV(N)$$

On peut effectuer un développement limité du potentiel au voisinage du point O en supposant que les dimensions du dipôle sont faibles devant la distance caractéristique de variation du champ. On a par définition du gradient :

$$dV = \overrightarrow{\text{grad}} V \cdot d\overrightarrow{OM}$$

soit :

$$V(M) - V(O) = \overrightarrow{\text{grad}}_M V(O) \cdot \overrightarrow{OM}$$

La notation $\overrightarrow{\text{grad}}_M$ signifie qu'on prend le gradient en considérant comme variables les coordonnées du point M et non celles du point O . La fonction est ensuite calculée au point O (il s'agit d'une généralisation à trois dimensions de la formule de Taylor : $f(x+h) = f(x) + hf'(x) + o(h)$).

En reportant ce résultat pour $M = N$ et $M = P$, on obtient :

$$\begin{aligned} Ep &= q \left(V(O) + \overrightarrow{\text{grad}}_M V(O) \cdot \overrightarrow{OP} - V(O) - \overrightarrow{\text{grad}}_M V(O) \cdot \overrightarrow{ON} \right) \\ &= q \overrightarrow{\text{grad}}_M V(O) \cdot \left(\overrightarrow{OP} - \overrightarrow{ON} \right) \\ &= q \left(-\overrightarrow{E}(O) \right) \cdot \overrightarrow{NP} \end{aligned}$$

D'où

$$Ep = -\overrightarrow{p} \cdot \overrightarrow{E}(O)$$

Cette énergie correspond d'après la définition de l'énergie électrostatique à l'énergie que doit fournir un opérateur extérieur pour amener le dipôle supposé rigide et déjà constitué depuis l'infini dans sa position actuelle.

On note que l'énergie potentielle d'interaction des deux charges constituant le dipôle reste constante dans le cas d'un dipôle rigide, ce qui explique qu'on s'intéresse surtout à l'énergie électrostatique dans un champ extérieur.

On peut également retrouver les résultats relatifs à l'équilibre d'un dipôle dans un champ électrostatique extérieur et à leur stabilité. En effet,

$$Ep = -\overrightarrow{p} \cdot \overrightarrow{E}(O) = pE \cos \theta$$

en notant θ l'angle entre le champ électrostatique et le dipôle. Le dipôle sera en équilibre lorsque l'énergie potentielle sera extrémale c'est-à-dire pour $\theta = 0$ ou $\theta = \pi$. Pour $\theta = 0$, $Ep = pE$ est maximale et pour $\theta = \pi$, $Ep = -pE$ est minimale. La position d'équilibre $\theta = 0$ est donc stable et la position $\theta = \pi$ instable ce qui est en accord avec l'étude du paragraphe 2.1.c).

b) Compléments : efforts exercés par un champ extérieur sur un dipôle rigide

De l'énergie potentielle d'interaction obtenue au paragraphe précédent, on peut déduire les efforts exercés sur le dipôle à savoir la résultante des forces et le moment résultant.

En effet, le travail élémentaire des forces et des moments s'écrit :

$$\begin{aligned}\delta W &= \vec{F} \cdot d\vec{OM} + \vec{\Gamma} \cdot d\vec{\Omega} \\ &= \vec{F} \cdot d\vec{OM} + \vec{\Gamma} \cdot d\theta \vec{k}\end{aligned}$$

en notant $d\vec{OM}$ le vecteur de translation élémentaire et $d\vec{\Omega} = d\theta \vec{k}$ le vecteur de rotation élémentaire.

On ne considère dans un premier temps qu'un déplacement en translation : $d\theta = 0$. Alors, par définition de l'énergie électrostatique,

$$\delta W = -\delta W_{\text{op}} = dEp$$

Donc

$$\vec{F} = -\overrightarrow{\text{grad}}(-\vec{p} \cdot \vec{E}) = \overrightarrow{\text{grad}}(\vec{p} \cdot \vec{E})$$

Pour le couple, on vérifie également qu'on retrouve l'expression générale à partir de l'énergie potentielle. On se place dans le plan perpendiculaire à $d\vec{\Omega}$ et on ne considère qu'un phénomène de rotation $d\vec{OM} = \vec{0}$. On utilise les coordonnées polaires dans le plan où a lieu la rotation. Alors, par analogie avec le gradient :

$$\Gamma_{\Omega} = -\frac{\partial Ep}{\partial \Omega}$$

Or, dans le cas d'une rotation du dipôle rigide, on a :

$$Ep = -\vec{p} \cdot \vec{E} = -pE \cos \Omega$$

D'où :

$$\Gamma_{\Omega} = -(-pE(-\sin \Omega)) = -pE \sin \Omega$$

Or la projection de $\vec{p} \wedge \vec{E}$ sur $\vec{u}_{\Omega}$ est

$$(\vec{p} \wedge \vec{E})_{\Omega} = -pE \sin \Omega$$

car Ω est l'angle entre $\vec{E}$ et $\vec{p}$.

On vérifie donc bien :

$$\vec{\Gamma} = \vec{p} \wedge \vec{E}$$

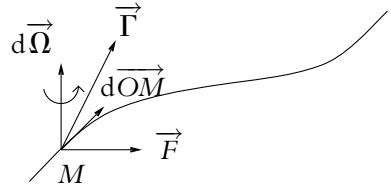


Figure 46.8 Efforts exercés sur un dipôle rigide dans un champ électrostatique extérieur.

c) Énergie potentielle d'un dipôle non rigide dans un champ extérieur

Dans la démonstration du paragraphe 2.3.a), on a intégré l'expression de dV . Or la position des points M à considérer (P ou N) dépend du champ électrique appliqué et on ne peut plus faire le calcul d'intégration aussi simplement. Il est alors nécessaire de revenir à la définition de l'énergie électrostatique.

Exemple

On suppose qu'un système présente un moment dipolaire $\vec{p} = \alpha \vec{E}$. On cherche à calculer l'expression de l'énergie potentielle d'interaction de ce dipôle avec un champ électrique extérieur nul par hypothèse à l'infini. On fait également l'hypothèse que $\vec{E} = E(x)\vec{u}_x$.

Le calcul de la force subie par le dipôle donne :

$$\vec{F} = \left(\vec{p} \left(\vec{E} \right) \cdot \left(\vec{\text{grad}} E \right) \right) \vec{u}_x = \left(\alpha E \frac{dE}{dx} \right) \vec{u}_x = \frac{\alpha}{2} \frac{dE^2}{dx} \vec{u}_x$$

L'énergie potentielle d'interaction s'obtient alors par :

$$Ep = W_{\text{op}} = - \int_{\infty}^x F_x dx = - \int_{\infty}^x \frac{\alpha}{2} dE^2 = - \frac{\alpha E^2}{2} = - \frac{\vec{p} \cdot \vec{E}}{2}$$

Cette expression est bien différente de celle obtenue pour un dipôle rigide :

$$Ep = - \vec{p} \cdot \vec{E}$$

3. Complément : approximation dipolaire

3.1 Position du problème

L'un des grands intérêts de l'étude du dipôle électrostatique réside dans le fait qu'il s'agit d'un modèle simple auquel nombre de distributions de charges peuvent se ramener à condition qu'on cherche à décrire leur action à grande distance de la distribution.

On va donc au cours de ce paragraphe considérer un ensemble de charges ponctuelles q_i placées en des points P_i , l'ensemble de ces charges se trouvant dans un volume fini de l'espace près d'un point O qu'on choisira comme origine du repère dans lequel on va se placer. L'objectif est de déterminer le potentiel puis le champ électrostatiques « loin » de cette distribution, à

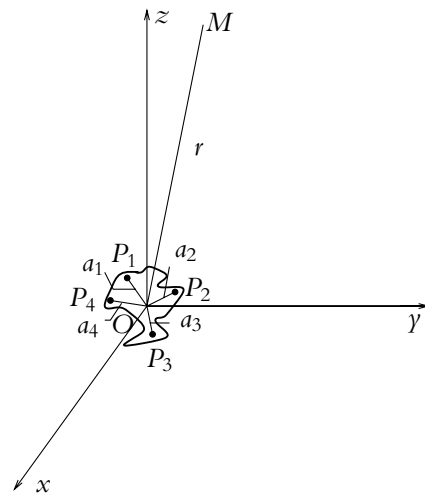


Figure 46.9 Distribution de charges et approximation dipolaire.

savoir en un point M de l'espace tel que :

$$\forall i, \|\overrightarrow{OP_i}\| \ll \|\overrightarrow{OM}\|$$

Cette condition permettra d'effectuer un développement limité des expressions en $\frac{a_i}{r}$ en notant $a_i = \|\overrightarrow{OP_i}\|$ et $r = \|\overrightarrow{OM}\|$.

3.2 Potentiel électrostatique créé par une distribution de charges ponctuelles en un point éloigné

On part de l'expression générale donnant le potentiel créé par une distribution de charges ponctuelles :

$$V(M) = \frac{1}{4\pi\epsilon_0} \sum_i \frac{q_i}{\|\overrightarrow{MP_i}\|}$$

Or :

$$\overrightarrow{MP_i} = \overrightarrow{MO} + \overrightarrow{OP_i}$$

D'où :

$$(MP_i)^2 = (MO)^2 + (OP_i)^2 + 2\overrightarrow{MO} \cdot \overrightarrow{OP_i}$$

Soit en introduisant les notations r et a_i définies plus haut :

$$\begin{aligned} \frac{1}{MP_i} &= \left(r^2 + a_i^2 - 2\overrightarrow{OM} \cdot \overrightarrow{OP_i} \right)^{-\frac{1}{2}} \\ &= \frac{1}{r} \left(1 + \frac{\overrightarrow{OM} \cdot \overrightarrow{OP_i}}{r^2} + \frac{3 \left(\overrightarrow{OM} \cdot \overrightarrow{OP_i} \right)^2 - a_i^2 r^2}{2r^4} + o \left(\frac{a_i}{r} \right)^2 \right) \end{aligned}$$

en se limitant à l'ordre 2 en $\frac{a_i}{r}$.

On en déduit :

$$V(M) = V_0(M) + V_1(M) + V_2(M) + \dots$$

en notant :

$$\left\{ \begin{array}{l} V_0(M) = \frac{1}{4\pi\epsilon_0} \frac{\sum_i q_i}{r} \\ V_1(M) = \frac{1}{4\pi\epsilon_0} \frac{\overrightarrow{OM} \cdot \left(\sum_i q_i \overrightarrow{OP_i} \right)}{r^3} \\ V_2(M) = \frac{1}{4\pi\epsilon_0} \frac{\sum_i \left(3q_i \left(\overrightarrow{OM} \cdot \overrightarrow{OP_i} \right)^2 - q_i a_i^2 r^2 \right)}{2r^5} \end{array} \right.$$

3.3 Distribution unipolaire

Si $\sum_i q_i \neq 0$ alors le terme prépondérant est $V_0(M)$ du fait que par hypothèse de travail : $\frac{a_i}{r} \ll 1$.

En choisissant comme origine O le barycentre des charges défini par :

$$\sum_i q_i \overrightarrow{OP_i} = \vec{0}$$

le terme $V_1(M)$ est nul. La distribution se comporte comme une charge ponctuelle de valeur $Q = \sum_i q_i$ placée en O :

$$V(M) = \frac{1}{4\pi\epsilon_0} \frac{Q}{r}$$

On dit qu'une telle distribution est *unipolaire*.

3.4 Distribution dipolaire

Dans le cas où $\sum_i q_i = 0$, le terme $V_0(M)$ est nul et le terme prépondérant devient $V_1(M)$.

Le barycentre des charges ne peut plus être défini car la somme totale des charges est nulle. En revanche, en séparant charges positives et charges négatives, on peut définir un barycentre P des charges positives et un barycentre N des charges négatives.

On note $q = \sum_{i, q_i > 0} q_i = \sum_{i, q_i < 0} (-q_i)$. Les positions de P et N sont alors données par :

$$\overrightarrow{ON} = \frac{\sum_{i, q_i < 0} (-q_i) \overrightarrow{OP_i}}{q} \quad \text{et} \quad \overrightarrow{OP} = \frac{\sum_{i, q_i > 0} q_i \overrightarrow{OP_i}}{q}$$

D'où

$$\sum_i q_i \overrightarrow{OP_i} = \sum_{i, q_i > 0} q_i \overrightarrow{OP_i} - \sum_{i, q_i < 0} (-q_i) \overrightarrow{OP_i} = q \overrightarrow{OP} - q \overrightarrow{ON} = q \overrightarrow{NP}$$

En notant $\vec{p} = q \overrightarrow{NP}$, on obtient

$$V_1(M) = \frac{1}{4\pi\epsilon_0} \frac{\overrightarrow{OM} \cdot \vec{p}}{r^3}$$

On est donc dans ce cas revenu à une distribution équivalente constituée de deux charges opposées placées aux barycentres des charges de même signe. Ceci correspond

à un dipôle de moment dipolaire :

$$\vec{p} = q\overrightarrow{NP}$$

On parle dans ce cas de *distribution dipolaire*.

3.5 Distribution quadripolaire

Pour certaines distributions, le moment dipolaire peut être nul : ce sera le cas lorsque le barycentre P des charges positives et N celui des charges négatives sont confondus.

Alors $\overrightarrow{NP} = \vec{0}$ et $\vec{p} = \vec{0}$.

Il faut regarder un ordre supplémentaire dans le développement multipolaire :

- $\sum_i q_i = 0$ donc $V_0(M) = 0$,
- $\vec{p} = \vec{0}$ donc $V_1(M) = 0$.

On considérera le terme

$$V_2(M) = \frac{1}{4\pi\epsilon_0} \frac{\sum_i \left(3q_i \left(\overrightarrow{OM} \cdot \overrightarrow{OP_i} \right)^2 - q_i a_i^2 r^2 \right)}{2r^5}$$

qui devient prépondérant. Il s'agit du terme *quadripolaire*, on parle de distribution quadripolaire ou quadrupolaire.

3.6 Application à l'étude des molécules

Un exemple où cette étude est particulièrement importante est constitué de l'analyse des molécules et de leurs interactions électrostatiques.

a) Moment dipolaire permanent

Les atomes dans leur état fondamental sont neutres et on peut en première approximation les considérer comme à symétrie sphérique. On en déduit que :

- la somme des charges est nulle :

$$\sum_i q_i = 0$$

- le barycentre des charges positives et celui des charges négatives sont confondus :

$$\sum_{i,q_i>0} q_i \overrightarrow{OP_i} = \sum_{i,q_i<0} (-q_i) \overrightarrow{OP_i}$$

On n'a donc pas de moment dipolaire :

$$\vec{p} = \vec{0}$$

On se retrouve dans le cas quadripolaire.

En revanche, les molécules, bien que globalement neutres, peuvent présenter un moment dipolaire non nul. En effet, la disposition spatiale des atomes et/ou leur différence de propriétés les rendent dissymétriques : les barycentres des charges positives et des charges négatives peuvent ne pas être confondus. Dans ce cas, on a une distribution dipolaire. On parle alors de *moment dipolaire permanent*.

On peut noter que les moments dipolaires permanents sont d'autant plus grands que la molécule est dissymétrique.

D'autre part, certaines molécules ne présentent pas de moment dipolaire mais possèdent un moment quadripolaire. C'est le cas par exemple du dioxyde de carbone CO_2 .

Il faudra tenir compte de ces propriétés intrinsèques aux molécules et de leurs conséquences lors de l'étude de l'interaction entre molécules.

Comme on l'a déjà signalé, on n'utilise pas l'unité du système international le coulomb mètre (C.m) : cette unité n'est pas adaptée aux ordres de grandeurs rencontrés en chimie. On préfère le debye de symbole D avec la règle de conversion :

$$1\text{D} = 3,336 \cdot 10^{-30} \text{C.m}$$

Quelques valeurs de moments dipolaires :

Molécule	Moment dipolaire permanent (D)
chlorure d'hydrogène HCl	1,08
eau H_2O	1,85
monoxyde de carbone CO	0,11
ammoniac NH_3	1,49
acide nitrique HNO_3	2,17
propène CH_3CHCH_2	0,35
éthanol $\text{CH}_3\text{CH}_2\text{OH}$	1,70
éthanal CH_3CHO	2,70

b) Moment dipolaire induit

Un atome ou une molécule ne possédant pas de moment dipolaire permanent peut en acquérir un sous l'action d'un champ électrostatique $\vec{E}$, par exemple à l'approche d'une molécule possédant un moment dipolaire permanent.

En effet, sous l'action d'un champ électrostatique, les charges positives et les charges négatives se déplacent en sens opposé : les barycentres des charges positives et des charges négatives font de même. Dans le cas d'une molécule sans moment dipolaire permanent, ils passent d'une situation où ils étaient confondus à une situation où ils ne le sont plus. Ce simple déplacement des barycentres des charges positives et des charges négatives permet aux molécules d'acquérir un moment dipolaire qu'on qualifie de *moment dipolaire induit*.

On notera que les moments dipolaires induits sont négligeables devant les moments dipolaires permanents (lorsque ceux-ci existent).

Dans l'approximation linéaire qui est valable pour les champs faibles (approximation qui peut donc se justifier dans le cas du champ créé par une molécule voisine par exemple), on a un moment dipolaire proportionnel au champ $\vec{E}$:

$$\vec{p} = \alpha \epsilon_0 \vec{E}$$

La molécule est *polarisable* et α est un coefficient appelé *polarisabilité*.

Il faudra également tenir compte dans l'étude des interactions entre molécules de ce type de moment dipolaire.

A. Applications directes du cours

1. Doublet électrostatique – Moment électrique p d'un dipôle, d'après CCP 2006

Ce problème est une question de cours. On conseille de le traiter sans consulter le cours et de revoir *a posteriori* les parties qui ont posé problèmes.

On considère un ensemble de n charges ponctuelles q_i , situées aux points S_i dans un volume fini $\mathcal{V}$, telles que $\sum_{i=1}^n q_i = 0$. On désigne par $\vec{p} = \sum_{i=1}^n q_i \vec{OS_i}$ le moment dipolaire de cette distribution, supposé non nul, O étant un point fixe appartenant à $\mathcal{V}$.

- Vérifier que l'expression du moment dipolaire de cette distribution est indépendante du choix de l'origine O .
- En déduire le moment dipolaire d'un doublet formé de deux charges ponctuelles $(-q)$ en S_1 et $(+q)$ en S_2 ($q > 0$).
- Dans la molécule HF, la distance entre le noyau d'hydrogène et le noyau de fluor vaut : $d = 0,92.10^{-10}$ m.

a) En première approximation, on suppose le caractère ionique de la liaison H-F avec transfert de l'électron de l'hydrogène sur l'atome de fluor. Cet électron étant associé à ceux du fluor, ils forment une sphère chargée négativement, centrée sur le noyau du fluor. Effectuer l'inventaire des charges (protons, électrons) présentes au niveau des noyaux d'hydrogène et de fluor dans la molécule HF. (Numéro atomique du fluor : 9)

b) Déterminer la valeur du moment dipolaire $\|\vec{p}\|$, en debye (D), de la molécule supposée à liaison ionique.

Données :

- charge élémentaire : $e = 1,6.10^{-19}$ C
- Debye : $1 \text{ D} = (1/3).10^{-29}$ C.m

c) En réalité, le moment dipolaire électrique expérimental de la molécule vaut 1,83 D. On désigne par H et F les positions des noyaux d'hydrogène et de fluor respectivement, et par G le barycentre des charges électroniques de la liaison H-F. En déduire la distance FG

4. Potentiel scalaire électrostatique $V(M)$

Les charges ponctuelles $(-q)$ et $(+q)$ d'un doublet sont placées respectivement aux points $S_1 (0, 0, -a/2)$ et $S_2 (0, 0, +a/2)$ du repère $(Oxyz)$ (figure 46.10).

On désigne par $p = \|\vec{p}\|$, le moment dipolaire du doublet, par M , un point courant de coordonnées sphériques (r, θ, φ) . Les vecteurs $(\vec{u}_r, \vec{u}_\theta, \vec{u}_\varphi)$ sont les vecteurs de base du système de coordonnées sphériques.

On pose $r_1 = \|\vec{S_1M}\|$, $r_2 = \|\vec{S_2M}\|$, $r = \|\vec{OM}\|$ et $\vec{r} = \vec{OM}$.

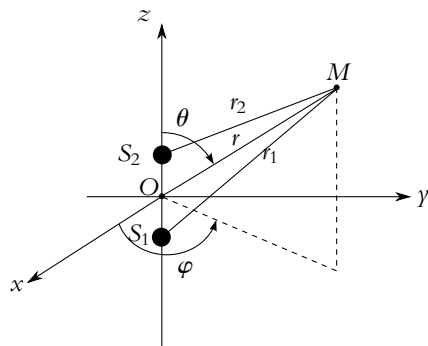


Figure 46.10

a) Exprimer le potentiel électrostatique $V(M)$ créé par le doublet, au point M , en fonction de q , r_1 et r_2 .

b) Établir son expression $V_d(M)$, pour un point M éloigné du doublet ($r \gg a$), en fonction de r , $\vec{r}$ et $\vec{p}$.

5. Champ électrostatique $\vec{E}(M)$

a) Montrer que $\vec{\text{grad}}_M(1/r^3)$ et $\vec{\text{grad}}_M(\vec{p} \cdot \vec{r})$ s'expriment en fonction de r , $\vec{r}$ ou $\vec{p}$ (l'indice M dans le gradient signifie que c'est le point M qui est la variable).

b) Dédire du potentiel $V_d(M)$ du dipôle, le champ électrostatique $\vec{E}(M)$ sous la forme :

$$\vec{E}(M) = \frac{1}{4\pi\epsilon_0} \left(\frac{k_l(\vec{p} \cdot \vec{r})\vec{r} - r^2\vec{p}}{r^5} \right)$$

où k_l est un facteur numérique que l'on calculera.

c) Déterminer les composantes (E_r , E_θ , E_φ) du champ $\vec{E}(M)$ en coordonnées sphériques.

d) La direction du champ en M est repérée par l'angle $\beta = (\vec{u}_r, \vec{E}(M))$. Quelle est alors la relation entre les angles β et θ ?

e) Calculer, dans le plan (yOz) limité au domaine $\theta \in [0, \pi/2]$ l'angle $\theta = \theta_1$ correspondant à un champ $\vec{E}(M)$ parallèle à l'axe Oy .

6. Équipotentielle et lignes de champ

a) Qu'appelle-t-on surfaces équipotentielles? Donner leur équation en coordonnées polaires pour ce dipôle.

b) Qu'appelle-t-on lignes de champ? Donner leur équation en coordonnées polaires.

c) Tracer, dans le plan (yOz) limité au domaine $\theta \in [0, \pi/2]$, l'allure de deux lignes équipotentielles ($V_1 > 0$ et $V_2 > V_1$) et de deux lignes de champ.

7. Action d'un champ électrique extérieur uniforme $\vec{E}_e$

On applique dans l'espace un champ extérieur $\vec{E}_e$.

a) Exprimer en fonction de $\vec{p}$ et de $\vec{E}_e$, le moment du couple $\vec{\Gamma}$ s'exerçant sur le dipôle.

b) L'énergie d'interaction U entre le dipôle et le champ extérieur $\vec{E}_e$ étant définie par : $U = -\vec{p} \cdot \vec{E}_e$, étudier les orientations d'équilibre du dipôle et préciser leur stabilité.

2. Potentiel créé par une répartition de charges quadrupolaire

Une répartition de charges est formée d'une charge $2q$ en un point O et de deux charges $-q$, sur l'axe Ox , l'une en $x = -a$ et l'autre en $x = a$. Cette distribution peut modéliser une molécule de dioxyde de carbone.

Calculer le potentiel créé à grande distance. En déduire le champ électrostatique correspondant.

3. Interaction de deux dipôles électrostatiques

On considère deux molécules possédant un moment dipolaire permanent qu'on assimile à deux dipôles rigides $\vec{p}_1$ et $\vec{p}_2$, leurs centres étant respectivement en P_1 et en P_2 . On s'intéresse dans cet exercice à leur interaction.

1. Exprimer l'énergie d'interaction de ces deux molécules.
2. On suppose dans la suite que les deux dipôles sont situés dans un même plan ; on note θ_1 l'angle entre $\vec{P_2P_1}$ et $\vec{p}_1$ et θ_2 l'angle entre $\vec{P_2P_1}$ et $\vec{p}_2$. Donner l'expression de l'énergie d'interaction en fonction de θ_1 et θ_2 .
3. On suppose que θ_1 est fixe. A quelle condition les deux molécules seront en équilibre l'une par rapport à l'autre ?
4. Expliciter la condition de stabilité de cet équilibre.
5. Donner pour les cas particuliers suivants les valeurs de θ_2 correspondant à un équilibre et préciser leur stabilité :
 - a) $\theta = 0$,
 - b) $\theta = \frac{\pi}{2}$,
 - c) $\theta = \pi$,
 - d) $\theta = -\frac{\pi}{2}$,
 - e) $\theta = \frac{\pi}{4}$.

4. Force de Van der Waals

Une molécule non polarisable D_0 placée au point O , a un moment dipolaire permanent $\vec{P}_0$. En un point A de l'axe du dipôle, à une distance r de O est placée une molécule D_1 dont le moment dipolaire $\vec{P}_1$ est induit par le champ électrostatique $\vec{E}$ créé par D_0 en A : $\vec{P}_1 = \alpha \epsilon_0 \vec{E}$ avec ($\alpha > 0$).

1. En appelant a la distance entre les deux charges du dipôle induit, calculer la force exercées par le champ $\vec{E}$ sur chacune d'elle.
2. En supposant la distance a petite devant r , déterminer la force s'exerçant sur D_1 .

B. Exercices et problèmes

1. Deux fils parallèles, d'après ENAC 2000

1. Un fil rectiligne infini de direction Oz porte des charges électriques positives uniformément réparties avec une densité linéique de charges λ . Déterminer le champ électrique créé par cette distribution en un point P situé à une distance $\rho = HP$ du fil. On note $\vec{u}_\rho$ et $\vec{u}_\varphi$ les vecteurs de la base polaire de P .

$$1. \vec{E}(P) = \frac{\lambda}{4\pi\epsilon_0\rho} \vec{u}_\rho,$$

$$2. \vec{E}(P) = \frac{\lambda}{2\pi\epsilon_0\rho} \vec{u}_\rho,$$

$$3. \vec{E}(P) = \frac{\lambda}{4\pi\epsilon_0\rho} \vec{u}_\varphi,$$

$$4. \vec{E}(P) = \frac{\lambda}{2\pi\epsilon_0\rho} \vec{u}_\varphi.$$

2. En déduire le potentiel $V(P)$ créé en P par le fil.

$$1. V(P) = -\frac{\lambda}{2\pi\epsilon_0} \ln \rho + \text{constante},$$

$$2. V(P) = \frac{\lambda}{4\pi\epsilon_0\rho^2} + \text{constante},$$

$$3. V(P) = \frac{\lambda}{4\pi\epsilon_0} \ln \rho,$$

$$4. V(P) = \frac{\lambda}{2\pi\epsilon_0} \rho (\ln \rho - 1).$$

3. On place deux fils du type précédent F_1 et F_2 l'un chargé positivement avec une densité de charges λ et l'autre chargé négativement avec une densité de charges $-\lambda$. Ils sont placés parallèlement à l'axe Oz à une distance d l'un de l'autre. Dans le plan xOy perpendiculaire à Oz , les fils laissent respectivement une trace en O_1 et O_2 sur l'axe Ox avec O milieu de $[O_1O_2]$. On considère un point M à une distance r de O , on note θ l'angle entre Ox et OM dans le plan xOy . La distance d est faible devant r , le produit λd restant constant et égal à p . Établir une expression approchée du potentiel $V(M)$ créé par les deux fils au point M . On utilisera les coordonnées polaires dans le plan xOy et on fixera l'origine des potentiels en O .

$$1. V(P) = \frac{p \cos^2 \theta}{4\pi\epsilon_0 r},$$

$$2. V(P) = \frac{p \sin \theta}{2\pi\epsilon_0 r},$$

$$3. V(P) = \frac{p \cos \theta}{2\pi\epsilon_0 r},$$

$$4. V(P) = \frac{p \sin \theta}{4\pi\epsilon_0 r}.$$

4. Déterminer la nature de la trace des équipotentiels dans le plan xOy autres que $V = 0$.

1. droites parallèles à l'axe Ox ,

2. droites parallèles à l'axe Oy ,

3. cercles centrés sur l'axe Oy ,

4. cercles centrés sur l'axe Ox .

5. Déterminer la nature des lignes sur lesquelles le module du champ est constant ainsi que l'angle α entre Ox et $\vec{E}(M)$.

1. ellipses dont l'un des foyers est O et $\alpha = \theta$,

2. droites parallèles à l'axe Ox et $\alpha = \frac{3\theta}{2}$,

3. cercles centrés en O et $\alpha = 2\theta$,

4. cercles centrés en O_1 ou O_2 et $\alpha = \frac{\theta}{2}$.

2. Modèle de solvation d'un ion

On peut représenter la molécule d'eau par un dipôle électrostatique de moment dipolaire $p = 6,2 \cdot 10^{-30}$ C.m.

On considère quatre molécules d'eau situées aux sommets G_1, G_2, G_3 et G_4 d'un tétraèdre régulier (figure 46.11). La distance séparant le centre du tétraèdre d'un sommet est $d = CG_i = 0,3$ nm.

On impose que l'axe des dipôles soit colinéaire au vecteur CG_i (mais pas forcément de même sens). Pour un couple de points (G_i, G_k) , l'angle $(G_iCG_k) = \beta = 109^\circ 28'$. On notera $k = 1/4\pi\epsilon_0 = 9 \cdot 10^9$ SI.

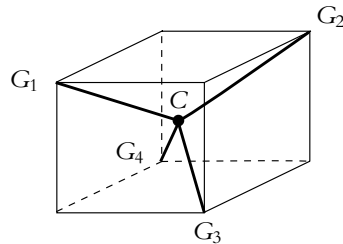


Figure 46.11

1. Établir l'expression de l'énergie d'interaction de 2 dipôles en fonction de p, d et β . Montrer que 3 cas peuvent se présenter.
2. En déduire les expressions de l'énergie d'interaction de ces quatre molécules d'eau en fonction de p, d et β pour les cinq arrangements des quatre dipôles respectant les conditions imposées.
3. On complète l'édifice en plaçant en C une charge ($e = 1,6 \cdot 10^{-19}$ C).
4. Calculer les nouvelles valeurs des énergies d'interaction correspondant aux cinq arrangements de la question 2. Application numérique : calculer cette énergie en $\text{kJ} \cdot \text{mol}^{-1}$ d'ions solvatés (nombre d'Avogadro $N_A = 6,02 \cdot 10^{23}$).
5. Quel est l'arrangement le plus stable des quatre molécules d'eau et de l'ion.
6. Que pensez-vous du modèle sachant que l'énergie de solvation d'un ion de cette taille est environ $-240 \text{ kJ} \cdot \text{mol}^{-1}$?

3. Cristal de NaCl soumis à un champ, d'après Polytechnique MP 2005

L'application d'un champ électrique à un monocristal de chlorure de sodium se traduit par des déplacements des ions qui le composent.

Soit $\vec{E} = E_0 \vec{u}_x$ le champ électrostatique imposé aux ions du cristal, par exemple en appliquant une tension à des électrodes planes plaquées sur les faces opposées d'un échantillon parallélépipédique. Sous l'effet de ce champ, les ions Na^+ se déplacent en bloc selon Ox de δ_+ et les ions Cl^- de δ_- , le centre de masse de l'ensemble restant immobile. On posera $x = \delta_+ - \delta_-$.

Avant tout déplacement le moment dipolaire global du cristal est nul par symétrie de la répartition. On note N le nombre d'ions sodium et chlorure par unité de volume et e la valeur absolue de la charge de l'électron.

1. Montrer que ces déplacements ioniques se traduisent par un moment dipolaire réparti, de densité volumique $\vec{P} = P\vec{u}_x$ et exprimer P en fonction de la charge élémentaire e , de x et du nombre N de paires d'ions Na^+Cl^- par unité de volume.
2. a) L'expérience montre que la relation entre P et E est linéaire, de la forme $P = \epsilon_0 \chi_{\text{ion}} E$ où χ_{ion} est un coefficient positif caractéristique du cristal. En déduire que le groupe d'ions

Na^+ est soumis à des forces de rappel élastique dont la moyenne par ion est de la forme $\overrightarrow{f} = -Kx\overrightarrow{u}_x$, les ions Cl^- étant soumis à des forces opposées.

b) Exprimer la constante K en fonction de N , e , ϵ_0 et χ_{ion} .

3. a) Après suppression du champ $\overrightarrow{E}$, les deux groupes d'ions évoluent librement. Écrire l'équation du mouvement d'un ion Na^+ et celle d'un ion Cl^- ; on désignera par m_+ et m_- leur masse respective.

b) Montrer que leur mouvement relatif est une oscillation à une pulsation ω_T que l'on explicitera en fonction de N , e , ϵ_0 , χ_{ion} et de m (masse réduite du couple d'ions).

47

Courant électrique et distributions de courants

Les sources du champ électrique sont des charges électriques fixes.

On s'intéresse maintenant aux sources du champ magnétique : les charges électriques en mouvement qui donnent lieu à des courants électriques. Pour suivre une approche analogue à celle de l'électrostatique, on commence l'étude de la magnétostatique par une description et une analyse des distributions de courants.

1. Courant électrique

1.1 Définition du courant électrique

On appelle *courant électrique* tout **mouvement d'ensemble** des particules chargées dans un référentiel.

Le point particulièrement important dans cette définition est le fait qu'il s'agisse d'un mouvement d'ensemble. En effet, toute particule est animée d'un mouvement descriptible au niveau microscopique dit d'agitation thermique, mouvement incessant et désordonné (Cf. cours de thermodynamique). Ce mouvement n'est pas forcément observé au niveau macroscopique, échelle pour laquelle on ne s'intéresse qu'au mouvement moyen d'un grand nombre de particules. À ce mouvement d'agitation thermique se superpose un mouvement d'ensemble de toutes les particules ; ce mouvement est lié à une action extérieure que toutes les particules subissent.

De ce fait, la vitesse d'une particule i peut être décomposée en la somme :

$$\vec{v}_i = \vec{v}_i^* + \vec{v}_{i,e}$$

où $\vec{v}_i^*$ est la vitesse liée à l'agitation thermique et $\vec{v}_{i,e}$ la vitesse due aux causes extérieures.

La vitesse moyenne de N particules est définie au niveau mésoscopique (sur un élément de volume $d\tau$ de l'ordre de $(10^{-6})^3 \text{ m}^3$ par exemple) par :

$$\langle \vec{v} \rangle = \frac{1}{N} \sum_{i=1}^N \vec{v}_i = \frac{1}{N} \sum_{i=1}^N \vec{v}_{i,e}$$

En effet, le mouvement d'agitation thermique est nul en moyenne car il s'agit d'un mouvement désordonné (Cf. cours de thermodynamique).

On s'intéresse donc ici au mouvement d'ensemble des particules chargées qui crée un courant électrique.

1.2 Divers types de courants électriques

On peut distinguer plusieurs types de courants électriques suivant leur origine.

a) Courant de conduction

Il s'agit du déplacement d'ensemble de particules chargées dans un milieu conducteur lié à l'existence d'un champ électrique $\vec{E}$.

Chaque particule de charge q , en plus de son mouvement d'agitation thermique, est soumise à une force :

$$\vec{F} = q \vec{E}$$

Toutes les charges de même type vont donc subir la même force et avoir un mouvement identique : on aura un mouvement d'ensemble donnant lieu à un courant électrique. L'origine électrique du courant permet de qualifier ce courant de *courant de conduction*.

b) Courant de convection

Les charges électriques sont parfois liées à des corps électriquement neutres et en mouvement. Le déplacement de ces corps entraîne celui des charges qui lui sont liées et l'existence d'un courant électrique liée à cet entraînement. On parle de *courant de convection*.

Dans ce type de courant, on classe aussi le mouvement des particules chargées qui se déplacent sous l'action d'un autre type de champ que le champ électrique (le champ de pesanteur par exemple).

Les courants de convection correspondent à l'entraînement des particules chargées sous une action autre qu'électromagnétique.

c) Courant de diffusion

Les *courants de diffusion* sont liés aux déplacements pouvant se produire du fait d'un gradient de concentration des particules chargées. Le mouvement tend à réduire cette différence.

Dans la suite, on ne s'intéressera pas à l'origine du courant ; le point important est l'existence d'un courant électrique quelle que soit sa nature.

1.3 Intensité électrique

On rappelle la définition donnée dans le cours d'électricité de première période pour l'intensité du courant électrique à travers une surface dS ; c'est la charge totale qui traverse cette surface par unité de temps :

$$I_{dS}(t) = \frac{\delta q}{dt}$$

en notant δq la charge traversant la surface dS entre t et $t + dt$.

2. Densité de courant

2.1 Définition de la densité de courant

On considère un fil conducteur où les porteurs de charges sont d'une seule et même nature (par exemple des électrons). On note ρ_m la densité volumique de charges mobiles.

Sous l'action d'un champ électrique extérieur, les porteurs de charges sont animés d'un mouvement d'ensemble à la vitesse moyenne $\vec{v}$.

On considère une surface $d\vec{S} = dS\vec{u}$ où $\vec{u}$ est un vecteur unitaire perpendiculaire à la surface dS . On cherche à déterminer la quantité de charges qui la traversent entre les instants t et $t + dt$.

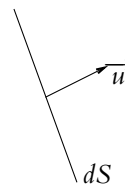


Figure 47.1 Surface élémentaire.

Entre t et $t + dt$, les porteurs de charges parcourent en moyenne vdt . Ceux qui traversent la surface S pendant cet intervalle de temps sont donc ceux contenus dans le volume correspondant en coupe à $ABCD$ (cf. figure 47.2) avec :

$$BC = vdt = AD$$

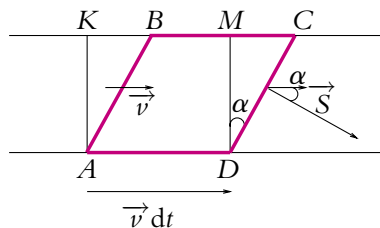


Figure 47.2 Densité de courant.

Ce volume est le même que celui qui correspond en coupe à $AKMD$ (les triangles AKB et DMC sont identiques). Le volume est donc le produit de la longueur vdt par la surface représentée en coupe par AK qui vaut : $dS \cos \alpha$. Le volume est donc :

$$v dS \cos \alpha dt = \vec{v} \cdot \vec{dS} dt$$

La charge qui traverse $\vec{dS}$ entre t et $t + dt$ est : $d^2q = \rho_m \vec{v} \cdot \vec{dS} dt$, ce qui permet d'écrire l'intensité élémentaire dI sous la forme :

$$dI = \rho_m \vec{v} \cdot \vec{dS} = \vec{j} \cdot \vec{dS}$$

en notant :

$$\vec{j} = \rho_m \vec{v}$$

Cette relation définit le *vecteur densité de courant* $\vec{j}$ dont la norme s'exprime en $A.m^{-2}$. L'intensité élémentaire dI est donc le flux du vecteur densité de courant à travers la surface orientée $d\vec{S}$.

► Remarques

1. Si $\vec{dS}$ et $\vec{v}$ sont colinéaires de même sens, $\alpha = 0$ et $dI = j dS$.
2. Si $\vec{j}$ n'est pas uniforme, $\vec{j}$ dépend de la position du point considéré. Il faut tenir compte de cette dépendance sur la surface, on calcule donc le flux de $\vec{j}$ à travers la surface S pour obtenir l'intensité du courant à travers S :

$$I = \iint_{M \in S} \vec{j}(M) \cdot \vec{dS}_M$$

3. S'il y a plusieurs types de charges ou de porteurs de charges, on généralise la définition donnée pour la densité de courant par :

$$\vec{j} = \sum_i \rho_{m,i} \vec{v}_i$$

la sommation portant sur les porteurs de charges. Ce sera notamment le cas des électrolytes comme par exemple une solution aqueuse de chlorure de sodium pour laquelle la densité de courant sera la somme de la densité de courant des ions chlorures Cl^- et de celle des ions sodium Na^+ :

$$\vec{j}_{sol} = \rho_{m,Cl^-} \vec{v}_{Cl^-} + \rho_{m,Na^+} \vec{v}_{Na^+}$$

La définition de l'intensité en fonction de $\vec{j}$ reste la même.

4. **Attention** : il ne faut pas confondre la densité volumique de charges mobiles ρ_m utilisée ici avec la densité volumique de charges ρ définie en électrostatique : toutes les charges ne sont pas en mouvement ! Par exemple, dans un conducteur, on a

un réseau cristallin fixe de densité de charges ρ_f et des éléments de conduction (par exemple des électrons) de densité de charges ρ_m . On ne s'intéresse ici qu'aux porteurs de charges mobiles. La neutralité électrique des conducteurs implique qu'en général :

$$\rho = \rho_f + \rho_m = 0$$

2.2 Lignes et tubes de courant

On définit :

- une *ligne de courant* comme une courbe en tout point tangente au vecteur densité de courant,
- un *tube de courant* comme l'ensemble des lignes de courant s'appuyant sur une courbe fermée.

2.3 Cas du régime stationnaire

En régime stationnaire, la densité volumique de courant ne varie pas au cours du temps. De plus, la conservation de la charge dans un volume d'un tube de courant délimité par deux sections S_1 et S_2 montre que le flux $\vec{j}$ à travers S_1 et à travers S_2 est le même (toutes les charges qui entrent en S_1 ressortent en S_2 et aucune charge ne sort par la surface latérale puisque $\vec{j}$ et $d\vec{S}$ sont perpendiculaires).

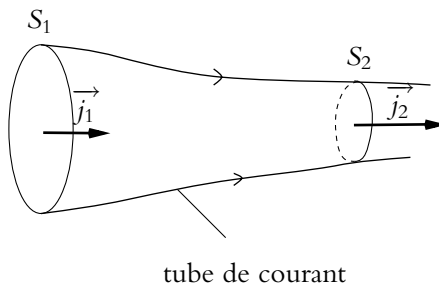


Figure 47.3 La densité de courant est à flux conservatif.

Cela impose à l'intensité d'être la même à travers toute section d'un même tube de courant.

On dit que $\vec{j}$ est à *flux conservatif*.

2.4 Différentes distributions de courant

a) Densité volumique de courant

La densité de courant qui vient d'être introduite est une densité volumique de courant au même titre qu'a été introduite en électrostatique une densité volumique de charges.

On rappelle que la densité volumique de courant s'écrit :

$$\vec{j} = \rho_m \vec{v}$$

où ρ_m est la densité volumique de charges **mobiles** et $\vec{v}$ la vitesse moyenne de ces charges mobiles.

La densité volumique de courant s'exprime en ampère par mètre carré (A.m^{-2}). Il est important de bien noter l'unité **par mètre carré**.

b) Densité surfacique de courant

Dans certains cas, les courants sont confinés au voisinage d'une surface S d'épaisseur e faible devant les autres dimensions du problème. Il est alors souvent souhaitable de considérer une densité surfacique de courant telle que :

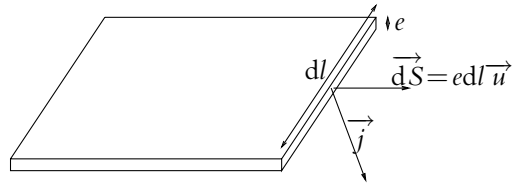


Figure 47.4 Densité surfacique de courants.

$$\vec{j}_{\text{surf}} = \lim_{e \rightarrow 0} (\vec{j} e)$$

ou encore $\vec{j}_{\text{surf}} = \int_0^e \vec{j} dl = \vec{j} e$ en intégrant sur l'épaisseur si on peut supposer que la densité de courant est constante sur l'épaisseur.

En effet, l'intensité du courant à travers la surface dS est :

$$dI = \vec{j} \cdot d\vec{S} = \vec{j} \cdot e dl \vec{u} = \vec{j}_{\text{surf}} \cdot dl \vec{u}$$

On en déduit l'expression pour e faible :

$$\vec{j}_{\text{surf}} = \vec{j} e$$

qui est l'expression au premier ordre de la densité surfacique de courant définie plus haut par une limite quand e tend vers 0. On note que le vecteur unitaire $\vec{u}$ est lié à la surface orientée et qu'il est donc perpendiculaire à dl .

La densité surfacique de courant j_{surf} s'exprime en A.m^{-1} .

Cette modélisation introduit des discontinuités qui peuvent être « résolues » en réintroduisant l'épaisseur de la surface comme cela a été vu en électrostatique à propos de la densité surfacique de charges.

c) Densité linéique de courant

Il existe également des cas où les courants sont localisés le long d'un fil qui est alors un tube de courant de faible section. Le volume élémentaire peut s'exprimer par :

$$d\tau = dl \cdot \vec{S}$$

On en déduit l'expression du courant élémentaire :

$$\vec{j} d\tau = \vec{j} (dl \cdot \vec{S})$$

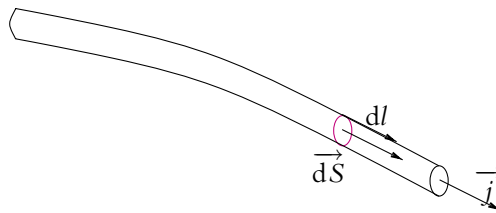


Figure 47.5 Densité linéique de courants.

Dans le cas d'un tube de courant, ces vecteurs $\vec{j}$, $d\vec{l}$ et $\vec{S}$ sont colinéaires donc on peut intervertir leurs positions relatives dans l'expression précédente :

$$\vec{j} d\tau = (\vec{j} \cdot \vec{S}) \cdot d\vec{l} = I d\vec{l}$$

soit

$$I = \vec{j} \cdot \vec{S}$$

Ceci correspond à une modélisation linéique qui sera très souvent utilisée. Ce sera la seule modélisation considérée dans la suite du cours de première année.

d) Cas d'une charge en mouvement

Pour une particule de charge q en mouvement à la vitesse $\vec{v}$, la densité de courant s'obtient directement à partir de la recherche de la quantité de charges qui traversent la surface dS entre les instants t et $t + dt$.

La quantité $q\vec{v}$ remplace ici $I d\vec{l}$.

On utilisera notamment ce résultat pour exprimer le courant lié au mouvement d'un électron autour du noyau décrit dans le cadre d'un modèle classique (l'électron décrit une orbite circulaire).

3. Invariances d'une distribution de courants

3.1 Invariance par translation

Une distribution de courants est invariante par translation lorsque l'image de la distribution par la translation est la distribution elle-même. Ceci n'est possible qu'avec une distribution s'étendant jusqu'à l'infini ; sinon la superposition entre la distribution et son image par translation n'est que partielle.

Ce type d'invariance permet de ne plus tenir compte de la dépendance selon cette direction. Par exemple, si on considère une distribution $\mathcal{D}$ invariante par toute translation le long de l'axe Ox alors $\vec{j}(x + a, y, z) = \vec{j}(x, y, z)$ pour toute valeur de a et donc $\vec{j}$ ne dépend pas de x :

$$\vec{j} = \vec{j}(y, z)$$

Ce sera par exemple le cas d'un fil rectiligne infini parcouru par une intensité I : on a invariance par toute translation le long de l'axe du fil.

3.2 Invariance par rotation

De même, l'invariance par rotation correspond au cas où la distribution obtenue après rotation se superpose rigoureusement avec la distribution initiale que ce soit en position dans l'espace ou en valeur locale de la densité de courants.

Ce type d'invariance permet de ne plus tenir compte de la dépendance selon l'angle de rotation considéré. Si on considère une distribution $\mathcal{D}$ invariante par toute rotation autour de l'axe Oz alors, en utilisant les coordonnées cylindriques d'axe Oz , $\|\vec{j}(r, \theta + \psi, z)\| = \|\vec{j}(r, \theta, z)\|$ pour toute valeur de ψ et donc la norme de $\vec{j}$ ne dépend pas de θ :

$$\|\vec{j}\| = \|\vec{j}(r, z)\|$$

Ce sera par exemple le cas d'un fil rectiligne parcouru par une intensité I : on a invariance par toute rotation autour de l'axe du fil.

Les invariances permettent donc de réduire le nombre de variables utiles.

48

Champ magnétique créé par des courants permanents

L'objet de ce chapitre est de définir le champ magnétique à partir de son action sur une ou plusieurs charges en mouvement. On pourra alors à partir des propriétés de symétrie des sources déduire les propriétés de symétrie du champ magnétique, énoncer la loi de Biot et Savart donnant l'expression du champ magnétique créé par une distribution de courants permanents et analyser l'ensemble des interactions d'origine magnétique.

1. Définition du champ magnétique

1.1 Observations expérimentales des actions magnétiques

En présence d'une source de champ magnétique (aimant par exemple), on constate les faits expérimentaux suivants concernant la force subie par une particule chargée :

- la particule ne subit pas de force pour une direction particulière de la vitesse notée $\vec{u}$,
- pour toute autre direction de la vitesse $\vec{v}$ de la particule, la force est perpendiculaire à la fois à $\vec{u}$ et à $\vec{v}$,
- l'intensité de la force est proportionnelle à :
 - la vitesse $\vec{v}$ de la particule,
 - au sinus de l'angle entre $\vec{v}$ et $\vec{u}$.

1.2 Force de Lorentz

Ces observations conduisent à écrire cette force, appelée *force de Lorentz*, sous la forme suivante :

$$\vec{F} = q \vec{v} \wedge \vec{B}$$

On définit ainsi le champ magnétique $\vec{B}$ par son action sur une particule de charge q et animée d'une vitesse $\vec{v}$ tout comme le champ électrostatique avait été défini par son action sur une particule de charge q au repos.

On notera que les champs électrique et magnétique dépendent du référentiel dans lequel on les considère (puisque la vitesse $\vec{v}$ en dépend).

1.3 Ordres de grandeur

Dans les unités du système international, la force s'exprime en newton (N), la charge en coulombs (C) et la vitesse en mètre par seconde (m.s⁻¹). L'unité du champ magnétique est donc parfaitement déterminée à partir de sa définition ; on l'appelle *tesla*, de symbole T.

Dans certains cas, on utilise une autre unité : le gauss de symbole G tel que

$$1\text{G} = 10^{-4}\text{T}$$

Quelques valeurs usuelles de champ magnétique :

Situation	Valeur (T)
champ magnétique terrestre	$0,2 \cdot 10^{-4} - 0,4 \cdot 10^{-4}$
entrefer d'un électroaimant	0,1 - 2
bobine supraconductrice	5 - 50

1.4 Principe de superposition

Soient deux champs magnétiques $\vec{B}_1$ et $\vec{B}_2$.

Une particule de charge q animée d'une vitesse $\vec{v}$ subit une force magnétique $\vec{F}_1$ due à $\vec{B}_1$ et une force magnétique $\vec{F}_2$ due à $\vec{B}_2$ avec :

$$\vec{F}_1 = q\vec{v} \wedge \vec{B}_1 \quad \text{et} \quad \vec{F}_2 = q\vec{v} \wedge \vec{B}_2$$

Au total, elle est soumise à l'action de la somme des forces comme le stipule le principe fondamental de la dynamique :

$$\vec{F} = \vec{F}_1 + \vec{F}_2 = q\vec{v} \wedge (\vec{B}_1 + \vec{B}_2) = q\vec{v} \wedge \vec{B}$$

Tout se passe comme s'il n'y avait qu'un seul champ magnétique :

$$\vec{B} = \vec{B}_1 + \vec{B}_2$$

Cela résulte de la linéarité des équations. Le principe de superposition s'applique au champ magnétique comme il s'appliquait au champ électrostatique.

2. Propriétés de symétrie du champ magnétique et conséquences

Seuls les résultats finaux de ce paragraphe sont à retenir. Les démonstrations proposées ne le sont qu'à titre d'approfondissement et sont l'adaptation au cas du champ magnétique de celles effectuées pour le champ électrostatique.

2.1 Invariances des lois de la magnétostatique

L'espace est supposé homogène isotrope : aucune direction et aucune origine n'est privilégiée.

De plus, la charge électrostatique est invariante par tout changement de référentiel et par tout déplacement.

Par conséquent, les lois de la magnétostatique doivent également respecter cette propriété d'invariance : quel que soit le déplacement envisagé, les lois de la magnétostatique sont les mêmes avant et après déplacement. Par exemple, la force de Lorentz s'exprime par la même relation avant et après.

Il s'agit de l'un des postulats de base de l'électromagnétisme : les lois de l'électromagnétisme sont invariantes par tout déplacement et par tout changement de référentiel¹ ; on a utilisé cette propriété en électrostatique, on va faire de même en magnétostatique.

2.2 Caractère axial du champ magnétique – Notion de pseudo-vecteur

On cherche la manière dont est transformé le champ magnétique par une symétrie par rapport à un plan.

Soit une charge ponctuelle q placée en M où règne un champ magnétique $\vec{B}$ et animée d'une vitesse $\vec{v}$. Elle subit la force :

$$\vec{F} = q \vec{v} \wedge \vec{B}$$

On considère la symétrie par rapport au plan $\mathcal{P}$. La charge q devient la charge $q' = q$ du fait de son caractère invariant et se retrouve en M' , symétrique de M par rapport à $\mathcal{P}$. Elle est

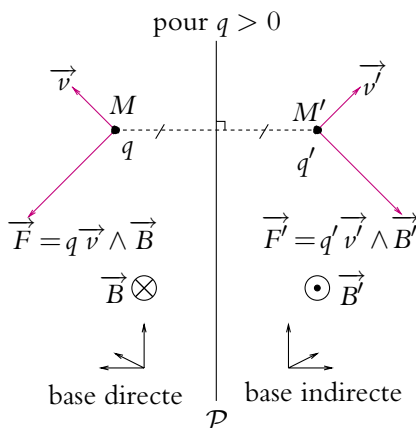


Figure 48.1 Champ magnétique et symétrie par rapport à un plan.

¹ C'est notamment le fait que ce postulat n'est pas toujours vérifié dans le cadre de la mécanique classique qui a conduit à élaborer la théorie de la relativité.

animée d'une vitesse $\vec{v}'$, symétrique de $\vec{v}$ par rapport à $\mathcal{P}$. Elle subit donc une force $\vec{F}'$ symétrique de $\vec{F}$ par rapport à $\mathcal{P}$.

L'invariance des lois de l'électromagnétisme permet d'avoir la même écriture pour la relation entre force, vitesse et champ magnétique :

$$\vec{F}' = q \vec{v}' \wedge \vec{B}'$$

La définition du produit vectoriel implique alors que :

$$\vec{B}' = -\vec{B} = -\text{sym}_{\mathcal{P}}(\vec{B})$$

L'image par la symétrie par rapport à un plan $\mathcal{P}$ du champ magnétique $\vec{B}$ est égal à l'opposé du symétrique de $\vec{B}$ par rapport au plan $\mathcal{P}$.

La démonstration qui vient d'être faite correspond au cas où le champ magnétique $\vec{B}$ est parallèle au plan $\mathcal{P}$ et perpendiculaire à la vitesse $\vec{v}$. Elle se généralise à toutes les positions relatives du champ magnétique $\vec{B}$.

Une telle propriété caractérise les pseudo-vecteurs ou vecteurs axiaux : le champ magnétique $\vec{B}$ est un *pseudo-vecteur* ou un *vecteur axial*. C'est dû au fait qu'il est défini par un produit vectoriel.

2.3 Distribution présentant un plan de symétrie

Soit une distribution $\mathcal{D}$ de densité de charges $\vec{j}$ admettant le plan $\mathcal{P}$ comme plan de symétrie.

$\mathcal{D}'$, la distribution symétrique de $\mathcal{D}$ par rapport au plan $\mathcal{P}$, est identique à la distribution $\mathcal{D}$ puisque $\mathcal{P}$ est un plan de symétrie. En notant $\vec{j}'$ sa densité de charges, on a donc pour tout point M de l'espace :

$$\vec{j}'(M) = \vec{j}(M') \quad (48.1)$$

où M' est le symétrique de M par rapport au plan $\mathcal{P}$.

D'autre part, on a établi au paragraphe précédent que :

$$\vec{B}'(M') = -\text{sym}_{\mathcal{P}}(\vec{B}(M)) \quad (48.2)$$

en notant $\vec{B}'$ le champ créé par la distribution $\mathcal{D}'$ symétrique de $\mathcal{D}$ par rapport au plan $\mathcal{P}$ et $\vec{B}$ le champ créé par la distribution $\mathcal{D}$.

Or d'après le principe de Curie, le champ magnétique créé en M' l'est soit par $\mathcal{D}$ soit par $\mathcal{D}'$, ce qui se traduit compte tenu (48.1) par

$$\vec{B}'(M') = \vec{B}(M')$$

En utilisant la relation (48.2), on obtient donc :

$$\vec{B}(M') = -\text{sym}_{\mathcal{P}} \left(\vec{B}(M) \right)$$

Le champ magnétique créé en un point M' symétrique d'un point M par rapport à un plan de symétrie de la distribution est égal à l'opposé du symétrique du champ magnétique créé en M .

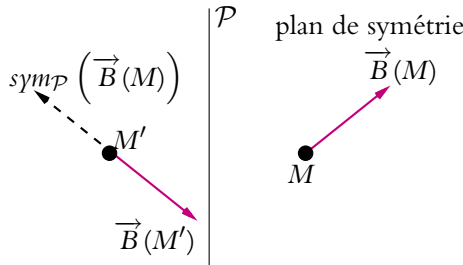


Figure 48.2 Champ magnétique et plan de symétrie.

Si M appartient à un plan de symétrie des sources ($\mathcal{P}$ par exemple), $M' = M$. On en déduit donc que :

$$\vec{B}(M) = -\text{sym}_{\mathcal{P}} \left(\vec{B}(M) \right)$$

ce qui se traduit en utilisant les indices t et n respectivement pour les composantes tangentielle et normale au plan $\mathcal{P}$:

$$\begin{cases} \vec{B}_t(M) = -\vec{B}_t(M) \\ \vec{B}_n(M) = \vec{B}_n(M) \end{cases}$$

d'où :

$$\vec{B}_t(M) = \vec{0}$$

Le champ magnétique n'a donc qu'une composante normale sur les plans de symétrie de la distribution qui le crée : il est perpendiculaire au plan de symétrie de la distribution de courants qui le crée.

Le champ magnétique est perpendiculaire aux plans de symétrie de la distribution qui le crée.

2.4 Distribution présentant un plan d'antisymétrie

Soit une distribution $\mathcal{D}$ de densité de charges $\vec{j}$ admettant le plan $\mathcal{P}$ comme plan d'antisymétrie.

$\mathcal{D}'$, la distribution symétrique de $\mathcal{D}$ par rapport au plan $\mathcal{P}$, est l'opposé de la distribution $\mathcal{D}$ puisque $\mathcal{P}$ est un plan d'antisymétrie. En notant $\vec{j}$ sa densité de charges, on a donc pour tout point M de l'espace :

$$\vec{j}(M) = -\vec{j}(M') \quad (48.3)$$

où M' est le symétrique de M par rapport au plan $\mathcal{P}$.

On a établi que :

$$\vec{B}'(M') = -\text{sym}_{\mathcal{P}} \left(\vec{B}(M) \right) \quad (48.4)$$

en notant $\vec{B}'$ le champ créé par la distribution $\mathcal{D}'$ symétrique de $\mathcal{D}$ par rapport au plan $\mathcal{P}$ et $\vec{B}$ le champ créé par la distribution $\mathcal{D}$.

Or d'après le principe de Curie, le champ magnétique créé en M' l'est soit par $\mathcal{D}$ soit par $\mathcal{D}'$, ce qui se traduit compte tenu de (48.3) par

$$\vec{B}'(M') = -\vec{B}(M')$$

En utilisant la relation (48.4), on obtient donc :

$$\vec{B}(M') = \text{sym}_{\mathcal{P}} \left(\vec{B}(M) \right)$$

Le champ magnétique créé en un point M' symétrique d'un point M par rapport à un plan d'antisymétrie de la distribution est égal au symétrique du champ magnétique créé en M .

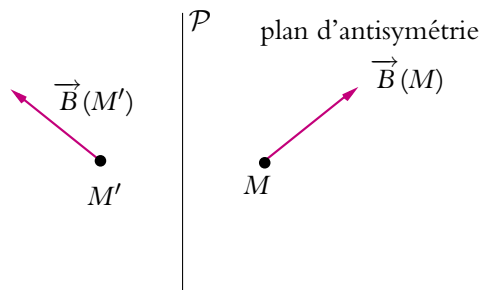


Figure 48.3 Champ magnétique et plan d'antisymétrie.

Si M appartient à un plan d'antisymétrie des sources (par exemple le plan $\mathcal{P}$), $M' = M$. On en déduit donc que :

$$\vec{B}(M) = \text{sym}_{\mathcal{P}} \left(\vec{B}(M) \right)$$

ce qui se traduit en utilisant les indices t et n respectivement pour les composantes tangentielle et normale au plan $\mathcal{P}$:

$$\begin{cases} \vec{B}_t(M) = \vec{B}_t(M) \\ \vec{B}_n(M) = -\vec{B}_n(M) \end{cases}$$

d'où :

$$\vec{B}_n(M) = \vec{0}$$

Le champ magnétique n'a donc que des composantes tangentielles sur les plans d'antisymétrie de la distribution qui le crée : il appartient au plan d'antisymétrie de la distribution de courants qui le crée.

Le champ magnétique appartient aux plans d'antisymétrie de la distribution qui le crée.

3. Loi de Biot et Savart

3.1 Énoncé de la loi de Biot et Savart

La loi de Biot et Savart a été postulée par ses auteurs sous la forme qui va être énoncée puisqu'elle permettait de rendre compte de la réalité du champ magnétique observé. On suivra la même démarche ici. Elle sera admise sans démonstration pour un circuit filiforme comme le stipule le programme.

Soit un circuit filiforme parcouru par un courant d'intensité I . Une longueur $d\vec{l}_P$ de ce circuit a une densité linéique de courant $I d\vec{l}_P$ dans le sens de parcours du courant.

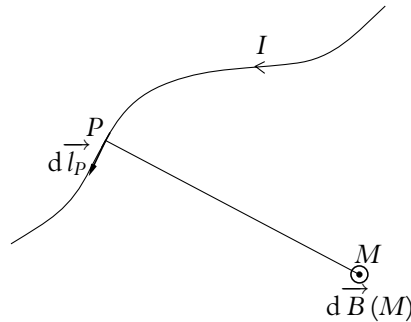


Figure 48.4

Le champ magnétique élémentaire créé par cet élément de courant s'écrit :

$$d\vec{B}(M) = \frac{\mu_0}{4\pi} I(P) d\vec{l}_P \wedge \frac{\vec{PM}}{(PM)^3}$$

Il ne s'agit que d'un intermédiaire de calcul correspondant à la contribution de l'élément de courant au champ magnétique total : il est en effet exclu d'isoler une portion

de circuit des autres. C'est une différence importante avec le champ électrostatique pour lequel il est parfaitement possible d'isoler une charge des autres. De ce fait, le champ magnétique élémentaire n'a pas de signification physique propre.

La quantité μ_0 est une constante universelle portant le nom de *perméabilité du vide*. Elle vaut :

$$\mu_0 = 4\pi 10^{-7} \text{ H.m}^{-1}$$

Elle est liée à la permittivité du vide ϵ_0 et à la vitesse de la lumière dans le vide c par la relation : $\epsilon_0 \mu_0 c^2 = 1$.

Le champ magnétique de la distribution s'obtient alors en intégrant sur les différents éléments de courant constituant la distribution :

$$\vec{B}(M) = \frac{\mu_0}{4\pi} \int_{P \in \text{circuit}} I(P) d\vec{l}_P \wedge \frac{\vec{PM}}{(PM)^3}$$

3.2 Complément : cas d'une charge en mouvement

Dans ce cas, l'élément de courant $I d\vec{l}_P$ s'écrit $q \vec{v}$ en notant q la valeur de la charge située au point P et $\vec{v}$ sa vitesse.

Dans le cas où $\vec{v}$ est quasi-constante dans l'espace où la charge se déplace et à la condition que sa valeur soit petite devant celle de la lumière, on obtient l'expression suivante pour le champ magnétique créé en un point M :

$$\vec{B}(M) = \frac{\mu_0}{4\pi} \frac{q \vec{v} \wedge \vec{PM}}{(PM)^3}$$

On vérifie sur cet exemple que les observations expérimentales sont en accord avec l'expression proposée par Biot et Savart : le champ magnétique $\vec{B}(M)$ est perpendiculaire à la vitesse de la particule.

Si les hypothèses précédentes ne sont pas vérifiées, on ne peut pas exprimer le champ magnétique $\vec{B}(M)$ créé au point M par cette expression : il faut considérer que les charges sont relativistes. D'autre part, les courants sont variables, ce qui sera exclu ici puisqu'on s'est placé en régime statique ou permanent.

On peut noter que l'approximation non relativiste qui n'est pas valable pour une charge seule le devient dans le cadre général où on considère un ensemble de courants permanents pour lequel on considère une vitesse moyenne.

3.3 Application : cas du fil rectiligne infini

L'objet de ce paragraphe n'est pas de traiter l'étude complète du champ magnétique créé par un fil rectiligne infini, cette étude sera faite au cours d'un chapitre ultérieur consacré aux calculs des champs magnétiques usuels. On ne s'intéresse qu'à

l'obtention d'un résultat utile par la suite, que l'expression de Biot et Savart permet d'avoir. Cette méthode n'est cependant pas la meilleure façon de traiter le problème et n'est donc pas à retenir en tant que telle.

$$\begin{aligned}\vec{B}(M) &= \frac{\mu_0}{4\pi} \int_{-\infty}^{+\infty} I d\vec{l} \wedge \frac{\vec{PM}}{(PM)^3} \\ &= \frac{\mu_0}{4\pi} \int_{-\infty}^{+\infty} Idz \vec{u}_z \wedge \frac{\vec{PM}}{(PM)^3}\end{aligned}$$

Or $\vec{PM} = r\vec{u}_r - z\vec{u}_z$

donc $\vec{u}_z \wedge \vec{PM} = r\vec{u}_\theta$ et en utilisant les notations de la figure ci-contre :

$$\vec{B}(M) = \frac{\mu_0}{4\pi} \int_{-\infty}^{+\infty} Idz \frac{r}{R^3} \vec{u}_\theta$$

On a deux variables non indépendantes z et R (r est une constante puisqu'elle repère le point où on calcule le champ et non les points de la distribution sur lesquels on intègre).

On va effectuer un changement de variables et utiliser φ comme nouvelle et unique variable. On notera qu'il s'agit d'un angle orienté comme indiqué sur le schéma (pour lequel $\varphi < 0$).

On a :

- $\cos \varphi = \frac{r}{R}$ donc $R = \frac{r}{\cos \varphi}$,
- $\tan \varphi = -\frac{z}{r}$ donc $z = -r \tan \varphi$ et $dz = -r(1 + \tan^2 \varphi) d\varphi = -\frac{r}{\cos^2 \varphi} d\varphi$,
- φ varie de $+\frac{\pi}{2}$ à $-\frac{\pi}{2}$ quand z varie de $-\infty$ à $+\infty$.

Donc

$$\begin{aligned}\vec{B}(M) &= \frac{\mu_0 I}{4\pi} \int_{+\frac{\pi}{2}}^{-\frac{\pi}{2}} \frac{-r}{\cos^2 \varphi} d\varphi \frac{r \cos^3 \varphi}{r^3} \vec{u}_\theta \\ &= \frac{\mu_0 I}{4\pi} \int_{-\frac{\pi}{2}}^{+\frac{\pi}{2}} \frac{\cos \varphi}{r} d\varphi \vec{u}_\theta = \frac{\mu_0 I}{4\pi r} [\sin \varphi]_{-\frac{\pi}{2}}^{+\frac{\pi}{2}} \vec{u}_\theta\end{aligned}$$

soit finalement :

$$\vec{B}(M) = \frac{\mu_0 I}{2\pi r} \vec{u}_\theta$$

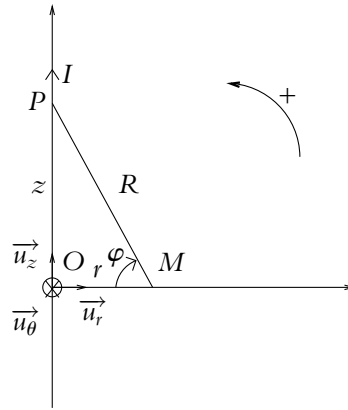


Figure 48.5 Fil rectiligne infini parcouru par un courant I .

4. Interactions magnétiques

4.1 Action d'un champ magnétique sur une particule chargée

La définition du champ magnétique est obtenue à partir de son action sur une particule de charge q animée d'une vitesse $\vec{v}$:

$$\vec{F} = q \vec{v} \wedge \vec{B}$$

C'est aussi à partir de cette force de Lorentz qu'on peut définir les interactions magnétiques.

4.2 Force de Laplace

Soit un circuit filiforme parcouru par une intensité I .

Les porteurs de charges de ce circuit subissent la force exprimée au paragraphe précédent soit, compte tenu de la distribution linéique de charges envisagée ici :

$$d\vec{F} = I d\vec{l} \wedge \vec{B}$$

On justifiera à partir de l'effet Hall² que cette force subie par les porteurs de charges d'un circuit filiforme se transmet à l'élément de circuit. Cette force dite *force de Laplace* :

$$d\vec{F} = I d\vec{l} \wedge \vec{B}$$

permet d'expliquer l'expérience du même nom. Une barre constituant un élément du circuit est parcourue par un courant d'intensité I . La présence d'un champ magnétique conduit à l'existence d'une force perpendiculaire à la barre, expliquant le mouvement de cette dernière.

4.3 Définition légale de l'ampère

Elle est basée sur l'interaction entre deux fils conducteurs infinis et parallèles.

D'après le calcul effectué plus haut, le fil 1 crée un champ magnétique :

$$\vec{B}_1(M) = \frac{\mu_0}{2\pi} \frac{I_1}{r} \vec{u}_\theta$$

en notant r la distance du point M au fil rectiligne et en utilisant les coordonnées cylindriques.

Une longueur l du fil 2 subit donc une force magnétique :

$$\vec{F}_{1 \rightarrow 2} = I_2 l \vec{u}_z \wedge \frac{\mu_0}{2\pi} \frac{I_1}{d} \vec{u}_\theta = -\frac{\mu_0}{2\pi} I_1 I_2 \frac{l}{d} \vec{u}_r$$

² Le phénomène de l'effet Hall sera étudié dans le chapitre sur le mouvement des particules chargées.

La connaissance du module de la force et des distances l et d permet de définir l'unité d'intensité, l'ampère, en faisant circuler le même courant dans les deux fils. On obtient la définition suivante :

L'ampère est l'intensité d'un courant constant qui, maintenu dans deux conducteurs rectilignes, infinis, parallèles, de section circulaire négligeable et distants de 1 m produit une force d'interaction entre ces deux conducteurs égale à $2 \cdot 10^{-7}$ N par mètre de conducteurs.

On fixe en même temps la constante μ_0 ou valeur de la perméabilité du vide à $4\pi 10^{-7}$ H.m⁻¹.

A. Applications directes du cours

1. Symétries de différentes distributions de courant

Déterminer les symétries des distributions suivantes :

1. un fil infini parcouru par un courant I ,
2. un fil de longueur L parcouru par un courant I ,
3. deux fils parallèles parcourus par des courants I_1 et I_2 , on étudiera le cas général puis les cas particuliers :
 - a) $I_1 = I_2$,
 - b) $I_1 = -I_2$,
4. une sphère bobinée c'est-à-dire sur laquelle on a enroulé un fil parcouru par un courant I ,
5. une nappe infinie plane de courants, parcourue par une densité surfacique de courants $\vec{j}_s = j_s \vec{u}_y$ uniforme,
6. une nappe plane de courants de largeur l suivant Oz et infinie dans la direction Oy , parcourue par une densité surfacique de courants $\vec{j}_s = j_s \vec{u}_y$ uniforme.

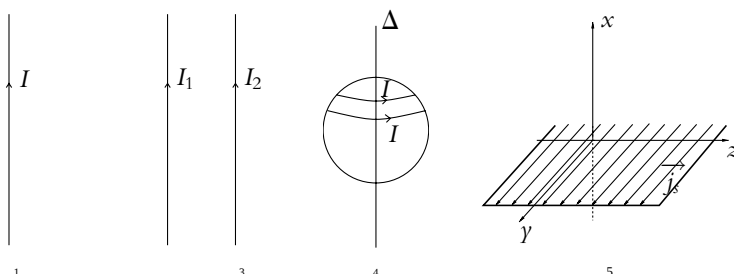


Figure 48.6

2. Interactions magnétiques entre différentes distributions

Décrire l'interaction entre les deux distributions dans les situations suivantes :

1. deux fils parallèles parcourus par des courants de même sens,
2. deux fils parallèles parcourus par des courants de sens contraire,
3. deux faisceaux d'électrons parallèles et de même vitesse,
4. deux faisceaux d'électrons parallèles et de vitesses opposées,
5. un fil parcouru par un courant et un faisceau d'électrons parallèle et de vitesse dans le même sens que le courant,
6. un fil parcouru par un courant et un faisceau d'électrons parallèle et de vitesse dans le sens opposé au courant.

B. Exercices et problèmes

1. Champ magnétique créé en un point de son axe par une demi-spire

On considère dans le plan xOy une demi-spire de centre O et de rayon a parcourue par un courant d'intensité I dans le sens des θ croissants si l'on repère un point P de la demi-spire en coordonnées cylindriques. L'angle θ est donc dans l'intervalle $\left[-\frac{\pi}{2}, \frac{\pi}{2}\right]$.

Soit M un point de l'axe Ox . Calculer la composante selon Oz du champ magnétique $\vec{B}$ en M .

2. Champ magnétique sur l'axe d'une hélice

On considère une hélice de rayon R et de pas h dont une équation paramétrique en coordonnées cartésiennes est : $x = R \cos \alpha$, $y = R \sin \alpha$ et $z = \frac{h\alpha}{2\pi}$. Cela peut constituer un modèle pour un solénoïde si l'hélice est parcourue par un courant d'intensité I .

1. Déterminer la composante selon Oz du champ magnétique en un point de l'axe (Oz). On pourra utiliser le changement de variables : $\alpha = \frac{2\pi R}{h} \cotan \varphi$.
2. Que devient ce résultat si on suppose l'hélice infinie ?
3. Déterminer l'expression de la composante axiale du champ magnétique faisant intervenir n , le nombre de spires par unité de longueur du solénoïde.
4. Quelle remarque peut-on formuler sur le résultat obtenu ?
5. Que se passe-t-il quand h tend vers 0 ?

Propriétés du champ magnétique

49

Ce chapitre est consacré à l'étude des propriétés du champ magnétique : flux, circulation et topographie.

1. Conservation du flux

1.1 Cas du fil rectiligne infini

On a établi au chapitre précédent à partir de la loi de Biot et Savart que le champ magnétique créé par un fil pouvait s'écrire en coordonnées cylindriques :

$$\vec{B} = \frac{\mu_0 I}{2\pi r} \vec{u}_\theta$$

Soit un tore de section S constante centré sur l'axe du fil parcouru par un courant.

Le flux du champ magnétique à travers toute section du tore est le même : le champ est colinéaire au vecteur surface orienté et la valeur du champ ne dépend que de la distance au fil qui est constante.

Le tore intercepte une surface fermée un nombre pair de fois, le flux étant alternativement « entrant » et « sortant ». Les contributions au flux sur la surface fermée sont donc opposées et au total s'annulent.

Pour décrire la totalité d'une surface fermée, il suffit de faire la même chose avec d'autres tores.

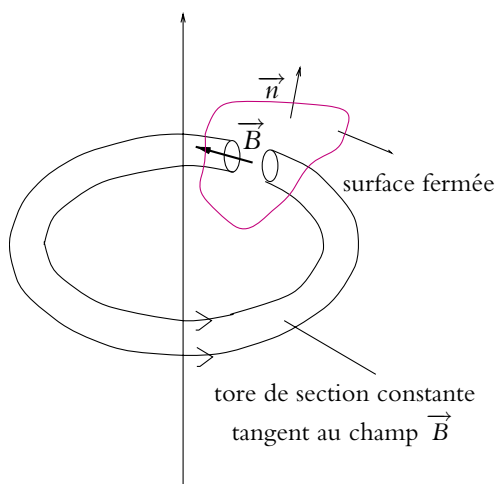


Figure 49.1 Conservation du flux du champ magnétique créé par un fil rectiligne infini.

Le flux du champ magnétique créé par un fil infini à travers une surface fermée est donc nul.

1.2 Cas d'un élément de courant

Soit un élément de courant $I \vec{dl} = Idl \vec{u}_z$ situé à l'origine O .

On a les mêmes propriétés de symétrie que pour le cas du fil rectiligne du paragraphe précédent. Il crée en $M(r, \theta, z)$ le champ magnétique élémentaire

$$d\vec{B}(M) = \frac{\mu_0}{4\pi} \frac{Idl \vec{u}_z \wedge (r \vec{u}_r + z \vec{u}_z)}{r^2 + z^2} = \frac{\mu_0 Idl}{4\pi} \frac{r}{r^2 + z^2} \vec{u}_\theta$$

On prend un tore identique à celui du paragraphe précédent et le flux du champ magnétique élémentaire sera le même à travers toute section du tore.

On obtient donc le même résultat au final : le flux du champ magnétique élémentaire à travers une surface fermée est nul.

1.3 Généralisation

D'après la loi de Biot et Savart, le champ magnétique dans le cas général résulte de la superposition des champs élémentaires :

$$\vec{B}(M) = \frac{\mu_0}{4\pi} \int_{P \in \mathcal{D}} I \vec{dl}_P \wedge \frac{\vec{PM}}{(PM)^3}$$

en notant $\mathcal{D}$ la distribution.

Le calcul du flux à travers une surface orientée est une opération linéaire :

$$\begin{aligned} \Phi &= \iint_{M \in \mathcal{S}} \vec{B}(M) \cdot d\vec{S}_M = \iint_{M \in \mathcal{S}} \left(\frac{\mu_0}{4\pi} \int_{P \in \mathcal{D}} I \vec{dl}_P \wedge \frac{\vec{PM}}{(PM)^3} \right) \cdot d\vec{S}_M \\ &= \iint_{M \in \mathcal{S}} \int_{P \in \mathcal{D}} d\vec{B}_P(M) \cdot d\vec{S}_M = \int_{P \in \mathcal{D}} \iint_{M \in \mathcal{S}} d\vec{B}_P(M) \cdot d\vec{S}_M = 0 \end{aligned}$$

car P et M sont indépendants l'un de l'autre, on peut donc intervertir l'ordre des intégrales.

On intègre sur la distribution des contributions de flux élémentaires nulles d'après le paragraphe précédent.

Le flux du champ magnétique à travers toute surface fermée est nul. On dit que le champ magnétique est à *flux conservatif*.

2. Circulation du champ magnétique Théorème d'Ampère

2.1 Exemple du fil rectiligne infini

On a établi que le champ créé par un fil rectiligne infini s'exprime en coordonnées cylindriques par :

$$\vec{B} = \frac{\mu_0 I}{2\pi r} \vec{u}_\theta$$

Le calcul de la circulation du champ magnétique le long d'un contour $\mathcal{C}$ donne en coordonnées cylindriques :

$$\begin{aligned} \oint_{M \in \mathcal{C}} \vec{B}(M) \cdot d\vec{OM} &= \oint_{\mathcal{C}} \frac{\mu_0 I}{2\pi r} \vec{u}_\theta \cdot (dr \vec{u}_r + r d\theta \vec{u}_\theta + dz \vec{u}_z) \\ &= \oint_{\mathcal{C}} \frac{\mu_0 I}{2\pi} d\theta = \frac{\mu_0 I}{2\pi} \oint_{\mathcal{C}} d\theta \end{aligned}$$

On doit distinguer le cas où le contour $\mathcal{C}$ encercle le fil rectiligne infini et le cas où il ne l'encercle pas :

- Le contour $\mathcal{C}$ encercle le fil rectiligne infini : on dit que $\mathcal{C}$ « enlace » le fil. Dans ce cas, on prend l'origine du repère à l'intérieur du fil et

$$\oint_{M \in \mathcal{C}} d\theta = \int_0^{2\pi} d\theta = 2\pi$$

On en déduit

$$\oint_{M \in \mathcal{C}} \vec{B}(M) \cdot d\vec{OM} = \mu_0 I$$

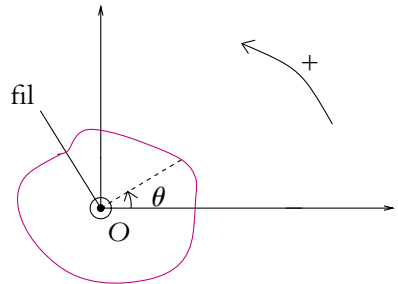


Figure 49.2 Fil encerclé par le contour.

- Le contour $\mathcal{C}$ n'encercle pas le fil rectiligne infini. On a alors la situation ci-contre.

$$\begin{aligned} \oint_{M \in \mathcal{C}} d\theta &= \int_{\theta(P_1)}^{\theta(P_2)} d\theta + \int_{\theta(P_2)}^{\theta(P_1)} d\theta \\ &= (\theta_2 - \theta_1) + (\theta_1 - \theta_2) = 0 \end{aligned}$$

On en déduit :

$$\oint_{M \in \mathcal{C}} \vec{B}(M) \cdot d\vec{OM} = 0$$

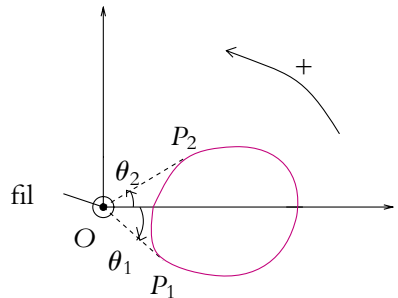


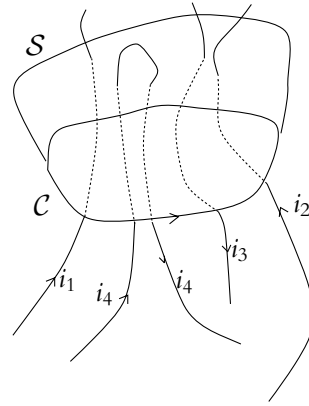
Figure 49.3 Fil à l'extérieur du contour.

2.2 Énoncé du théorème d'Ampère

Ce résultat se généralise à toute distribution de courants par le théorème d'Ampère.

La circulation du champ magnétique le long d'un contour fermé $\mathcal{C}$ est égale au produit de μ_0 par l'intensité totale qui traverse une surface quelconque¹ s'appuyant sur $\mathcal{C}$:

$$\oint_{M \in \mathcal{C}} \vec{B}(M) \cdot d\vec{OM} = \mu_0 I_{\text{enlacé}}$$



Sur l'exemple de la figure ci-contre,

$$I_{\text{enlacé}} = i_1 + i_2 - i_3 + i_4 - i_5 = i_1 + i_2 - i_3$$

Figure 49.4 Théorème d'Ampère.

Ce résultat est admis ici et sera établi dans le cours de deuxième année à partir des équations de Maxwell.

3. Topographie

Comme pour le champ électrostatique, on étudie la topographie du champ magnétique et on trace les cartes de champs. On définit donc les lignes de champ et les tubes de champ comme pour n'importe quel champ vectoriel.

3.1 Lignes de champ

On rappelle qu'une ligne de champ est définie comme la courbe constamment tangente au champ, ici au champ magnétique.

On l'obtient en exprimant la colinéarité du vecteur déplacement élémentaire et du champ magnétique à savoir :

$$\vec{B}(M) \wedge d\vec{OM} = \vec{0} \quad \text{ou} \quad \vec{B}(M) = k(M) \vec{OM}$$

où $k(M)$ est un réel dépendant de M .

La ligne de champ se réduit à un point dans le cas où le champ est nul.

¹Le résultat est indépendant du choix de la surface car la densité de courants est à flux conservatif en régime stationnaire.

3.2 Propriétés des lignes de champ

Au paragraphe 1, on a vu que localement le champ magnétique a la même structure que le champ magnétique créé par un fil rectiligne infini. Il est donc orthoradial et « tourne » autour des courants qui le créent. Les lignes de champs du champ magnétique s'enroulent donc autour des sources qui le créent.

Ceci constitue une différence essentielle avec le champ électrique. Les lignes de champ du champ électrique convergent ou divergent vers les charges électriques ou sources qui le créent.

Cette propriété peut être représentée par les figures suivantes :

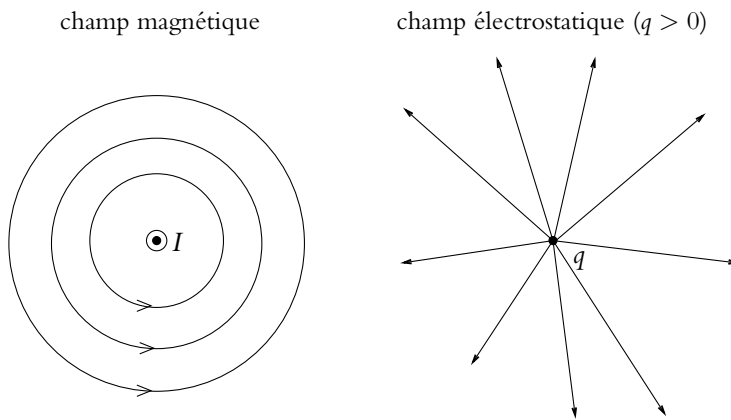


Figure 49.5 Allures des lignes de champ au voisinage de la source

Cette propriété des lignes de champ permet de distinguer les champs électrostatique et magnétique à partir du calcul de leur flux à travers une surface fermée ou du calcul de leur circulation le long d'un contour fermé.

3.3 Conséquence sur le flux

Dans le cas du champ électrique, les lignes de champ sont radiales par rapport à une charge ponctuelle : elles traversent donc un nombre impair de fois toute surface fermée entourant la charge et le flux du champ électrique à travers cette surface est non nul. La relation entre flux et charge s'obtient par le théorème de Gauss.

Au contraire, les lignes de champ magnétique sont orthoradiales et s'enroulent autour des courants qui le créent : elles peuvent donc être parallèles à la surface fermée et sinon, elles coupent la surface un nombre pair de fois, ce qui conduit à une contribution nulle au flux du champ magnétique à travers cette surface fermée. Le champ magnétique est à flux conservatif.

3.4 Conséquence sur la circulation

Dans le cas du champ électrique, les lignes de champ sont radiales par rapport à une charge ponctuelle : elles peuvent donc être perpendiculaires au contour fermé, ce qui conduit à une contribution nulle à la circulation du champ électrique sur le contour fermé. Le champ électrostatique est à circulation conservative.

Au contraire, les lignes de champ magnétique sont orthoradiales et s'enroulent autour des courants qui le créent : elles peuvent donc être parallèles au contour fermé, ce qui conduit à une contribution non nulle à la circulation du champ magnétique sur le contour fermé. La relation entre circulation et courant s'obtient par le théorème d'Ampère.

3.5 Exemples de topographie d'un champ

a) Deux fils infinis parallèles parcourus par un courant de même intensité dans le même sens

Soit la distribution de courants constituée de deux fils infinis parallèles parcourus par un courant d'intensité I circulant dans le même sens. Du fait de l'invariance par translation parallèlement à l'axe des fils, on peut limiter l'étude à un plan perpendiculaire aux fils. En utilisant les coordonnées cartésiennes, on suppose que les fils coupent ce plan aux points de coordonnées $(-1, 0)$ et $(1, 0)$. Un logiciel de simulation permet d'obtenir la carte des lignes de champ qui est donnée par les figures 49.6.

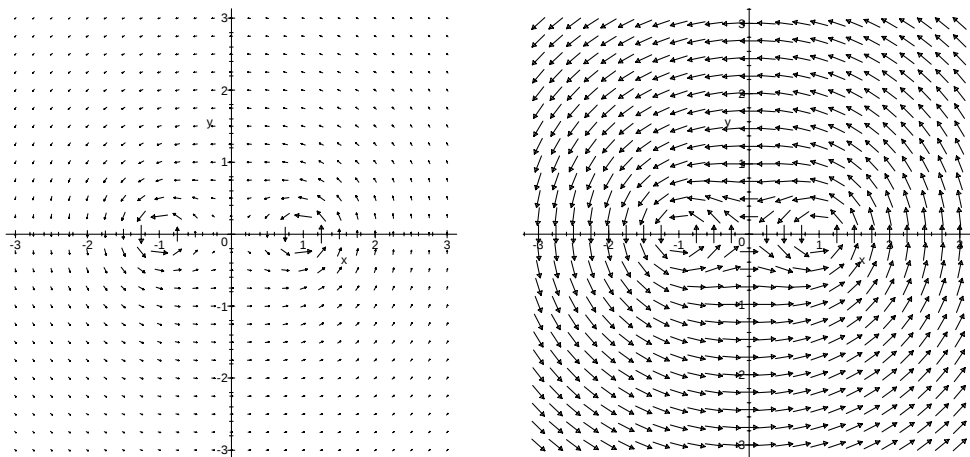


Figure 49.6 Lignes de champ de la distribution de courants constituée de deux fils parallèles parcourus par un courant de même intensité dans le même sens. La longueur des flèches est proportionnelle à la norme du champ sur la figure de gauche ; sur la figure de droite, les flèches ont même longueur indépendamment de la norme du champ.

On constate que les lignes de champ s'enroulent autour des fils dans le même sens de rotation : le courant circule dans le même sens.

La norme du champ diminue lorsqu'on s'éloigne des fils : le champ est plus important à proximité des sources qui le créent comme c'est le cas également pour le champ électrostatique.

On note l'existence d'un point d'arrêt où le champ est nul à mi-chemin entre les deux fils.

D'autre part, loin de la distribution, on a la même carte de champ qu'avec un seul fil.

Enfin, les plans $x = 0$ et $y = 0$ sont des plans de symétrie de la distribution de courant. Le champ magnétique est orthogonal à ces plans en un point de ces plans et, pour deux points M et M' symétriques par rapport à l'un de ces plans $\mathcal{P}$, on a $\vec{B}(M') = -\text{sym}_{\mathcal{P}}(\vec{B}(M))$.

b) Deux fils infinis parallèles parcourus par un courant de même intensité dans des sens opposés

On considère la même distribution que précédemment en inversant le sens du courant dans l'un des fils. Avec les mêmes notations que précédemment, on obtient les cartes de champ données par les figures 49.7.

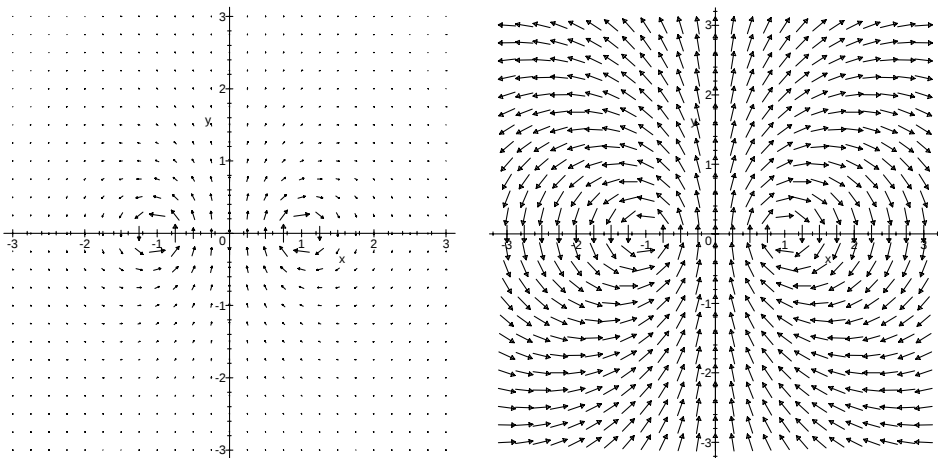


Figure 49.7 Lignes de champ de la distribution de courants constituée de deux fils parallèles parcourus par un courant de même intensité dans des sens opposés. La longueur des flèches est proportionnelle à la norme du champ sur la figure de gauche ; sur la figure de droite, les flèches ont même longueur indépendamment de la norme du champ.

On peut formuler les mêmes remarques que précédemment. Les différences sont liées au sens du courant qui n'est pas ici le même dans les deux fils. Les lignes de champ s'enroulent donc dans deux sens différents et il n'y a pas de point d'arrêt où le champ

magnétique s'annule. De plus, loin de la distribution, on ne retrouve pas la même carte des lignes de champ qu'avec un seul fil, contrairement à l'exemple précédent.

La norme du champ décroît également lorsqu'on s'éloigne des fils.

Enfin, le plan $x = 0$ est plan de symétrie de la distribution de courant. Le champ magnétique est orthogonal à ce plan en un point de ce plan et, pour deux points M et M' symétriques par rapport ce plan $\mathcal{P}$, on a $\vec{B}(M') = -\text{sym}_{\mathcal{P}}(\vec{B}(M))$. De plus, le plan $y = 0$ est plan d'antisymétrie de la distribution de courant. Le champ magnétique est tangent à ce plan en un point de ce plan et, pour deux points M et M' symétriques par rapport ce plan $\mathcal{P}$, on a $\vec{B}(M') = \text{sym}_{\mathcal{P}}(\vec{B}(M))$.

A. Applications directes du cours

1. Intersection de lignes de champ magnétique

Les lignes de champ d'un champ magnétique peuvent-elles se couper ?

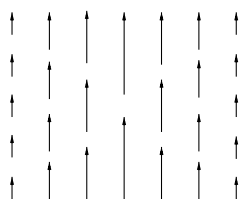
2. Allure des lignes de champ magnétique de quelques distributions

Tracer les lignes du champ magnétiques dans un plan perpendiculaire au(x) fil(s) dans les situations suivantes :

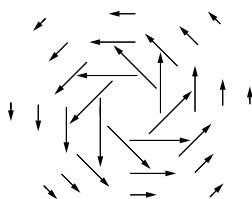
- deux fils parallèles distants de d et parcourus par des courants de même intensité et de même sens,
- deux fils parallèles distants de d et parcourus par des courants de même intensité et de sens opposé,
- trois fils parallèles aux sommets d'un triangle équilatéral et parcourus par des courants de même intensité et de même sens,
- la superposition d'un fil parcouru par un courant d'intensité I et d'un champ magnétique uniforme perpendiculaire au fil.

3. Lignes de champ

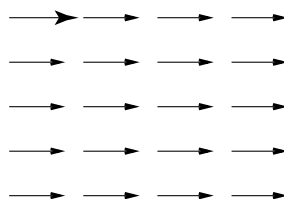
On donne les lignes de champ suivantes :



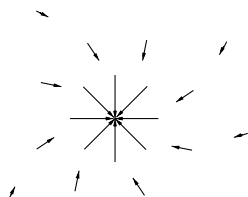
(a)



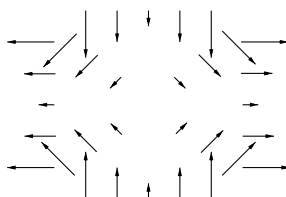
(b)



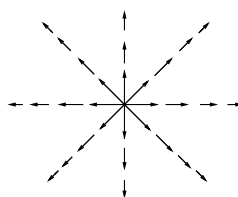
(c)



(d)



(e)



(f)

Préciser celles qui peuvent correspondre aux lignes de champ d'un champ magnétique. Si oui, proposer une distribution pouvant les avoir engendrées.

B. Exercices et problèmes

1. Solénoïde semi-infini

On considère un solénoïde S infiniment fin, de rayon a , d'axe (Oz) , comportant n spires par unité de longueur, et parcouru par un courant d'intensité I . Il est semi-infini, il s'étend le long du demi-axe $z < 0$. On appelle O le centre de la face terminale. Le champ créé par un solénoïde infini d'axe Oz est $\vec{B}_{\text{int}} = \mu_0 n I \vec{u}_z$ à l'intérieur et $\vec{B}_{\text{ext}} = \vec{0}$ à l'extérieur.

1. Déterminer l'expression de la composante suivant (Oz) du champ magnétique B créé par le solénoïde en un point M de sa face terminale.

2. Soit $\mathcal{L}$ une ligne de champ magnétique. Elle coupe la face terminale en un point P situé à la distance γ de O . À l'intérieur du solénoïde, elle tend à se confondre avec une parallèle à l'axe (Oz) .

a) Justifier cette affirmation.

b) Soit x la distance à l'axe de cette parallèle. Déterminer la relation entre x et γ .

2. Circulation du champ magnétique

Calculer la circulation du champ magnétique créé par une spire circulaire parcourue par un courant d'intensité I le long de son axe. On pourra faire un changement de variables du type :

$$u = \frac{z}{\sqrt{z^2 + a^2}}.$$

Que peut-on en conclure ?

L'expression du champ magnétique créé par une spire circulaire le long de son axe est établie dans le chapitre suivant ou dans l'exercice 3.

3. Champ magnétique au voisinage de l'axe d'une spire, d'après ENSTIM 2002

On donne une spire circulaire de rayon R , de centre O , d'axe Oz . Cette spire est parcourue par un courant électrique d'intensité I constante.

1. a) Montrer par des arguments de symétrie que, sur l'axe, le champ magnétostatique $\vec{B}$ est porté par l'axe et prend la forme de $\vec{B} = B(z)\vec{u}_z$ où $\vec{u}_z$ est un vecteur unitaire porté par l'axe Oz .

b) Comparer $B(z)$ et $B(-z)$.

c) Calculer le champ magnétostatique créé en un point M de l'axe tel que $OM = z$. On écrira $B(z) = B_0 f\left(\frac{z}{R}\right)$ où $B_0 = B(0)$. Préciser B_0 et $f\left(\frac{z}{R}\right)$.

d) Tracer le graphe de la fonction $B(z)$.

2. On s'intéresse maintenant au champ magnétostatique au voisinage de l'axe. On calcule donc le champ en un point M défini par des coordonnées cylindriques (r, θ, z) .

- a) Montrer par des arguments de symétrie très précis, qu'en M , le champ $\vec{B}$ n'a pas de composante orthoradiale B_θ .
- b) Montrer que la norme de $\vec{B}$ ne dépend que de r et z .
- c) Que peut-on dire du flux de $\vec{B}$ à travers une surface fermée ?
- d) Montrer que la circulation de $\vec{B}$ au voisinage de l'axe est conservative.
- e) Calculer le flux de $\vec{B}$ à travers une surface fermée cylindrique d'axe Oz dont les bases sont des disques de rayon r petit et de cotes z et $z + dz$. En déduire :

$$B_r(r, z) = -\frac{r}{2} \frac{dB_z(z, 0)}{dz}$$

Calculer l'expression de $B_r(r, z)$.

- f) Calculer de même la circulation du champ magnétique le long du petit rectangle de sommets $M(r, z)$, $P(r, z + dz)$, $Q(r + dr, z + dz)$ et $R(r + dr, z)$, en supposant r petit et dr et dz encore plus petits. En déduire que :

$$\frac{\partial B_z}{\partial r} = \frac{\partial B_r}{\partial z}$$

puis que :

$$B_z(r, z) = B_z(0, z) - \frac{r^2}{4} \frac{d^2 B_z}{dz^2}(0, z)$$

- g) On se place au voisinage du centre O d'une spire de rayon R . De combien peut-on s'écarter dans le plan de la spire pour que la composante axiale du champ magnétique diffère du champ B_0 au centre de moins de 1% ?

On donne : $\frac{d^2 B}{dz^2}(z) = -\frac{3B_0}{R^2} \left(1 - 4\frac{z^2}{R^2}\right) \left(1 + \frac{z^2}{R^2}\right)^{-7/2}$.

50

Exemples de champs magnétiques

On applique les différents résultats obtenus au cours des chapitres précédents pour déterminer l'expression du champ magnétique créé par quelques distributions classiques de courants.

1. Méthodes de calculs

1.1 Démarche générale

Le calcul d'un champ magnétique s'obtient en effectuant successivement les étapes suivantes¹ :

- Étude des invariances :
Cela permet de choisir le système de coordonnées adapté au problème : cartésiennes ou cylindriques en cas d'invariance par translation et cylindriques ou sphériques en cas d'invariance par rotation. Les invariances fournissent également les variables dont dépend le champ.
- Analyse des symétries :
La connaissance des plans de symétrie et d'antisymétrie permet de déterminer l'orientation du champ magnétique en un point : le champ magnétique appartient aux plans d'antisymétrie et est perpendiculaire aux plans de symétrie de la distribution de courants.
- Choix d'une méthode de calcul parmi les deux exposées dans la suite. Elles ne sont qu'au nombre de deux et non trois comme pour le champ électrostatique puisqu'on n'a pas défini de quantité analogue au potentiel électrostatique.

¹On notera que les deux premières étapes sont interchangeables. De plus, on remarquera qu'il s'agit de la même démarche que pour la détermination du champ électrostatique.

1.2 Calcul direct par la loi de Biot et Savart

Il s'agit d'intégrer la loi de Biot et Savart sur la distribution envisagée.

A priori, il est nécessaire de calculer trois intégrales correspondant aux composantes du champ sur une base fixe. Il ne faut pas oublier alors que l'étude des symétries peut avoir fourni des informations sur la direction du champ et donc que certaines composantes de celui-ci sont nulles. On évitera dans ce cas d'effectuer le calcul pour arriver dans le meilleur des cas (c'est-à-dire si aucune erreur de calcul n'a été commise !) à une valeur nulle.

1.3 Utilisation du théorème d'Ampère

Le théorème d'Ampère joue un rôle analogue pour le champ magnétique à celui que joue le théorème de Gauss pour le champ électrostatique. Il ne sera utilisable que lorsque l'étude des symétries et des invariances aura fourni la direction du champ, limité le nombre de variables dont il dépend et que seule manquera l'information relative à sa norme. Dans ce cas, il sera un outil précieux et performant.

Après avoir déterminé les invariances et les symétries et après avoir opté pour cette méthode, il convient de réaliser les deux étapes suivantes.

- Détermination du contour d'Ampère :
Celui-ci est défini de manière à rendre simple le calcul de la circulation. On choisira donc un contour composé de parties soit tangentes au champs le long desquelles la norme de $\vec{B}$ est uniforme, soit perpendiculaires au champ. Dans le premier cas, la circulation le long de cette partie du contour ne fait pas intervenir d'angle (les deux vecteurs sont colinéaires) et le calcul de la circulation sera très simple puisque la norme du champ est uniforme sur cette partie du contour. Dans le second, la circulation est nulle : le vecteur déplacement élémentaire et le champ magnétique sont perpendiculaires, leur produit scalaire est nul. Il est d'autre part nécessaire de bien tenir compte du fait suivant : **le contour d'Ampère (comme c'était le cas pour la surface de Gauss) doit passer par le point où on cherche à calculer le champ.**
- Calcul :
Il consiste à appliquer le théorème d'Ampère compte tenu des résultats antérieurs.

2. Fil rectiligne infini

Soit un fil infini parcouru par un courant I . On cherche le champ magnétique créé par ce fil en tout point de l'espace. Cet exemple, explicitement au programme, a déjà été traité par le calcul direct ; ici on le reprend en utilisant les propriétés du champ.

2.1 Étude des invariances

Le problème privilégie une seule direction, celle du fil : on choisit les coordonnées cylindriques en prenant l'axe du fil pour axe Oz .

Le fil est infini : la distribution est invariante par translation suivant l'axe Oz .

$\vec{B}$ ne dépend donc pas de z .

Elle est également invariante par rotation autour de l'axe Oz donc $\|\vec{B}\|$ ne dépend pas non plus de θ .

Les invariances de la distribution impliquent donc :

$$\|\vec{B}(M)\| = B(r)$$

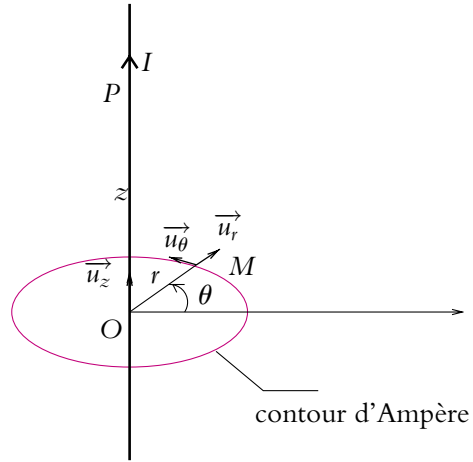


Figure 50.1 Fil rectiligne infini.

2.2 Étude des symétries

Le plan passant par M et contenant l'axe Oz est un plan de symétrie de la distribution. Le champ magnétique en M est donc perpendiculaire à ce plan : il est orthoradial.



On note que la recherche des plans de symétrie est plus intéressante dans le cas du champ magnétique que celle des plans d'antisymétrie : on a directement l'orientation du champ.

Par exemple, le plan passant par M et perpendiculaire à l'axe Oz est un plan d'antisymétrie, mais cette information ne donne pas la direction du champ magnétique en M : il appartient à ce plan. Au final

$$\vec{B}(M) = B(r)\vec{u}_\theta$$

2.3 Calcul de la norme du champ magnétique

On n'a plus qu'à déterminer la norme du champ magnétique : on va donc appliquer le théorème d'Ampère. Ce dernier donne la valeur de la circulation du champ magnétique le long d'un contour fermé $\mathcal{C}$:

$$\oint_{M \in \mathcal{C}} \vec{B}(M) \cdot d\vec{OM} = \mu_0 I_{\text{enlacé}}$$

Il faut donc choisir judicieusement le contour fermé : on le prend à une distance r de l'axe du fil afin que la valeur du champ soit constante sur le contour et que le champ

soit colinéaire à l'élément $d\overrightarrow{OM}$. Ce sera donc le cercle centré sur le fil, de rayon r , passant par le point M .

Dans ce cas, on a : $d\overrightarrow{OM} = r d\theta \overrightarrow{u_\theta}$ et

$$\oint_{M \in C} \overrightarrow{B}(M) \cdot d\overrightarrow{OM} = \int_0^{2\pi} B(r) r d\theta = 2\pi r B(r) = \mu_0 I_{\text{enlacé}} = \mu_0 I$$

On en déduit :

$$\overrightarrow{B}(M) = \frac{\mu_0 I}{2\pi r} \overrightarrow{u_\theta} = \frac{\mu_0 I}{2\pi} \overrightarrow{u_z} \wedge \frac{\overrightarrow{OM}}{r^2}$$

2.4 Topographie

Le champ est orthoradial donc les lignes de champ sont des cercles dont le centre est situé sur l'axe Oz .

Leur sens est donné par l'orientation du courant dans le fil.

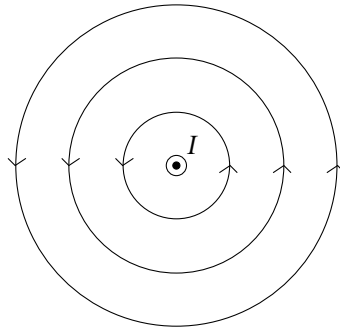


Figure 50.2 Ligne du champ magnétique créé par un fil infini dans le plan perpendiculaire au fil.

2.5 Complément : cas où on ne néglige plus la section du fil

Le résultat obtenu pose un problème pour déterminer le champ magnétique sur l'axe du fil : quand r tend vers 0, le champ magnétique devient infini. On va donc modifier la modélisation de la distribution en considérant une section a non nulle du fil.

On suppose dans ce cas que l'on a une densité volumique de courant

$$\overrightarrow{j} = j \overrightarrow{u_z} = \frac{I}{\pi a^2} \overrightarrow{u_z}$$

pour $r < a$. En effet, par définition de $\overrightarrow{j}$, on a :

$$I = \iint_{\text{section}} \overrightarrow{j} \cdot d\overrightarrow{S} = \iint_{\text{section}} j dS = j \iint_{\text{section}} dS = j \pi a^2$$

en appliquant successivement le fait que $\vec{j}$ et $d\vec{S}$ sont colinéaires et que la densité de courant est uniforme sur la section.

Les invariances et les symétries sont les mêmes que lorsque la section est négligeable donc

$$\vec{B}(M) = B(r)\vec{u}_\theta$$

On applique le théorème d'Ampère au même contour :

$$2\pi r B(r) = \mu_0 I_{\text{enlacé}}$$

La différence provient du calcul de $I_{\text{enlacé}}$ pour lequel on distingue :

- Si $r < a$, $I_{\text{enlacé}} = \pi r^2 j$ et :

$$\vec{B}(M) = \frac{\mu_0 r j}{2} \vec{u}_\theta = \frac{\mu_0 r I}{2\pi a^2} \vec{u}_\theta$$

On peut écrire ce résultat sous la forme :

$$\vec{B}(M) = \frac{\mu_0}{2} \vec{j} \wedge \overrightarrow{OM}$$

en prenant O un point de l'axe. Cette formulation est intéressante par son indépendance vis-à-vis du système de coordonnées utilisé et elle sera notamment utile lorsqu'on devra sommer les contributions de plusieurs cylindres.

- Si $r > a$, $I_{\text{enlacé}} = I = \pi a^2 j$ et

$$\vec{B}(M) = \frac{\mu_0 I}{2\pi r} \vec{u}_\theta$$

On peut également écrire ce résultat sous la forme :

$$\vec{B}(M) = \frac{\mu_0 a^2}{2} \vec{j} \wedge \frac{\overrightarrow{OM}}{r^2}$$

en prenant O un point de l'axe.

On retrouve le même résultat à l'extérieur du fil que lorsque la section était négligée.

On note que le champ magnétique est continu en $r = a$.

3. Spire circulaire

Soit une spire de centre O , d'axe Oz et de rayon R .

3.1 Symétries et invariances pour un point quelconque

a) Étude des invariances

L'axe de la spire est une direction privilégiée du problème : on choisit donc les coordonnées cylindriques.

La distribution n'est pas invariante par translation le long de l'axe. En revanche, elle est invariante par toute rotation autour de l'axe de la spire : le champ magnétique est indépendant de θ .

b) Étude des symétries

Le plan passant par M et contenant l'axe de la spire est un plan d'antisymétrie des sources (les courants tournent dans deux sens différents avant et après symétrie par rapport au plan).

$\vec{B}$ appartient à chacun de ces plans donc $\vec{B}$ a une composante radiale et une composante axiale.

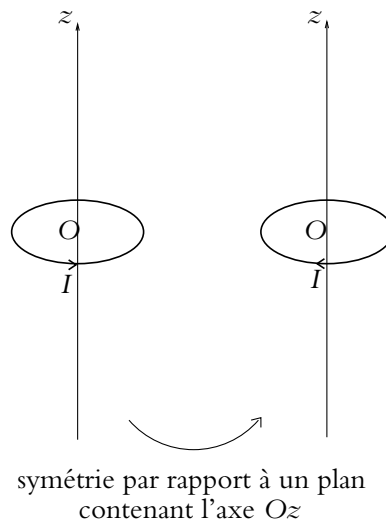


Figure 50.3 Spire circulaire et symétrie par rapport à un plan.

c) Limitation du problème

En un point quelconque, les symétries de la distribution ne sont donc pas suffisantes pour simplifier l'expression donnée par la loi de Biot et Savart.

Il faut recourir à une intégration numérique pour calculer la valeur du champ magnétique, ce qui est toujours théoriquement possible.

Dans la suite, on se limite au calcul du champ magnétique sur l'axe de la spire circulaire.

3.2 Symétries et invariances pour un point de l'axe

Comme on se place sur l'axe, la seule variable du problème est z donc :

$$\vec{B} = \vec{B}(z)$$

Tout plan contenant l'axe est plan d'antisymétrie donc le champ magnétique appartient à chacun de ces plans et est dirigé le long de l'axe Oz .

Au final :

$$\vec{B}(M \in Oz) = B(z)\vec{u}_z$$

Le plan $\mathcal{P}$ de la spire étant un plan de symétrie de la distribution de courant et $\vec{B}$ porté par $\vec{u}_z$, $\vec{B}(M') = \vec{B}(M)$ pour M' symétrique de M par rapport à $\mathcal{P}$ ou encore $B(z) = B(-z)$. Dans la suite, on pourra donc limiter les calculs au cas $z > 0$.

3.3 Calcul de la valeur du champ par la loi de Biot et Savart

$$\vec{B}(M) = \frac{\mu_0}{4\pi} \oint_{P \in \text{spire}} I d\vec{l}_P \wedge \frac{\vec{PM}}{(PM)^3}$$

Or la seule composante non nulle d'après l'étude des invariances est la composante axiale : on n'effectue l'intégration que pour B_z :

$$\begin{aligned} B_z &= \frac{\mu_0}{4\pi} \oint_{P \in \text{spire}} I \left(d\vec{l}_P \wedge \frac{\vec{PM}}{(PM)^3} \right) \cdot \vec{u}_z \\ &= \frac{\mu_0}{4\pi} \oint_{P \in \text{spire}} I \left(\frac{\vec{PM}}{(PM)^3} \wedge \vec{u}_z \right) \cdot d\vec{l}_P \\ &= \frac{\mu_0 I}{4\pi} \oint_{\text{spire}} \frac{\sin \alpha}{(PM)^2} dl_P \end{aligned}$$

car $\vec{PM} \wedge \vec{u}_z = PM \sin \alpha \vec{u}_\theta$ parallèle à $d\vec{l}_P$.

On peut retrouver ce résultat à partir des projections :

$$d\vec{l}_P = dl \vec{u}_\theta \quad \text{et} \quad \vec{PM} = -R\vec{u}_r + z\vec{u}_z$$

donc

$$d\vec{l}_P \wedge \vec{PM} = -R dl \vec{u}_\theta \wedge \vec{u}_r + z dl \vec{u}_\theta \wedge \vec{u}_z = R dl \vec{u}_z + z dl \vec{u}_r$$

On en déduit :

$$B_z = \frac{\mu_0 I}{4\pi} \oint_{P \in \text{spire}} \frac{R dl}{(PM)^3}$$

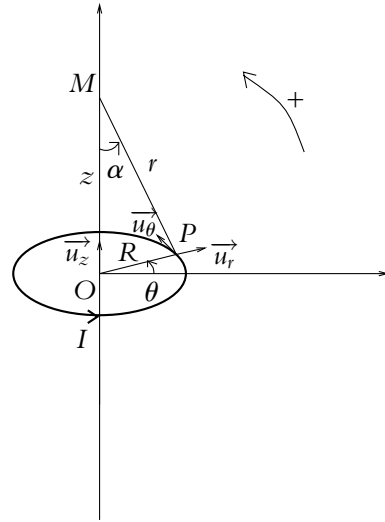


Figure 50.4 Champ magnétique créé par une spire circulaire.

Alors comme $PM = r = \frac{R}{\sin \alpha}$ prend la même valeur pour tout point de la spire et $dl = R d\theta$, on intègre sur θ entre 0 et 2π soit :

$$B_z = \frac{\mu_0 I \sin^3 \alpha}{4\pi R} \int_0^{2\pi} d\theta = \frac{\mu_0 I \sin^3 \alpha}{2R} \quad \text{soit} \quad \vec{B}(M) = \frac{\mu_0 I \sin^3 \alpha}{2R} \vec{u}_z$$

En remarquant que $\sin \alpha = \frac{R}{\sqrt{R^2 + z^2}}$, on peut aussi l'exprimer en fonction de z :

$$\vec{B}(M) = \frac{\mu_0 I}{2} \frac{R^2}{(R^2 + z^2)^{\frac{3}{2}}} \vec{u}_z$$

3.4 Topographie du champ magnétique créé par la spire

Elle s'obtient par intégration numérique à partir de la loi de Biot et Savart.

L'expression de champ magnétique au voisinage de l'axe Oz a été calculée à l'exercice B3 du chapitre précédent.

On a l'allure ci-dessous en se plaçant dans un plan perpendiculaire à la spire le long d'un de ses diamètres.

On retrouve le fait que les lignes de champ tournent autour des sources (courants) qui le créent dans le sens donné par exemple par la règle du tire-bouchon. C'est tout à fait différent du champ électrostatique pour lequel les lignes de champ divergent des sources (charges).

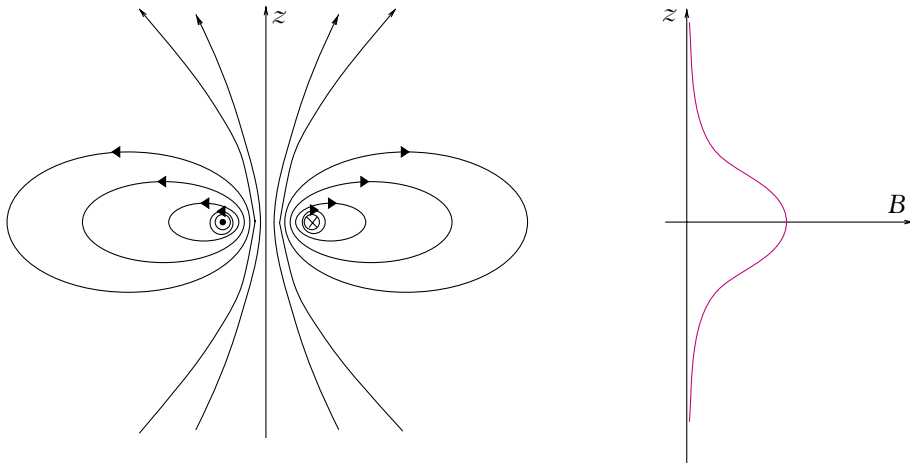


Figure 50.5 Topographie du champ magnétique créé par une spire circulaire.

3.5 Application aux bobines de Helmholtz

Une application usuelle du calcul du champ créé par une spire est l'étude des bobines de Helmholtz. Ce dispositif permet d'obtenir un champ quasi-uniforme le long de l'axe entre les deux bobines.

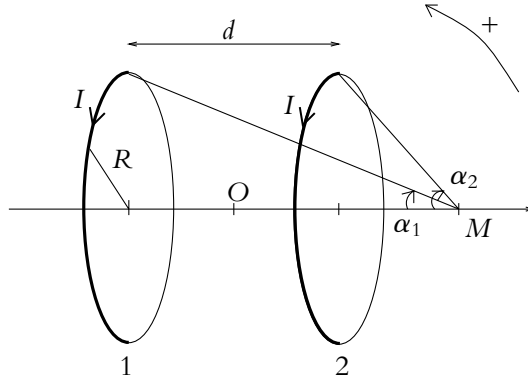


Figure 50.6 Bobines de Helmholtz.

Il s'agit de deux spires identiques parallèles parcourues par un même courant circulant dans le même sens, ces deux spires étant éloignées l'une de l'autre d'une distance d .

On cherche le champ en un point M de l'axe des deux spires.

On fixe l'origine sur l'axe à égale distance des deux spires.

D'après le principe de superposition, le champ est la somme des champs créés par chacune des deux spires :

$$\vec{B} = \vec{B}_1 + \vec{B}_2$$

Or

$$\vec{B}_i = \frac{\mu_0 I}{2R} \sin^3 \alpha_i \vec{u}_z$$

avec

$$\sin \alpha_1 = \frac{R}{\sqrt{R^2 + \left(z + \frac{d}{2}\right)^2}} \quad \text{et} \quad \sin \alpha_2 = \frac{R}{\sqrt{R^2 + \left(z - \frac{d}{2}\right)^2}}$$

En reportant ces expressions dans celles du champ, on obtient :

$$\vec{B} = \frac{\mu_0 I}{2R} \left(\left(1 + \frac{\left(z + \frac{d}{2}\right)^2}{R^2} \right)^{-\frac{3}{2}} + \left(1 + \frac{\left(z - \frac{d}{2}\right)^2}{R^2} \right)^{-\frac{3}{2}} \right) \vec{u}_z$$

On cherche la position relative des deux spires pour que le champ au voisinage du point O soit quasiuniforme : on veut donc des variations, représentées par $\frac{dB}{dz}$, minimales soit :

$$\frac{d^2B}{dz^2}(0) = 0$$

Un logiciel de calcul formel ou une calculatrice peut aider au calcul de cette double dérivation et permettre d'obtenir :

$$\frac{d^2B}{dz^2}(0) = \frac{3\mu_0 I}{R^3} \left(1 + \left(\frac{d}{2R} \right)^2 \right)^{-\frac{7}{2}} \left(\frac{d^2}{R^2} - 1 \right)$$

Donc :

$$\frac{d^2B}{dz^2}(0) = 0 \Leftrightarrow d = R$$

Dans ce cas, le champ aura l'expression suivante :

$$\vec{B} = \frac{\mu_0 I}{2R} \left(\left(1 + \left(\frac{z + \frac{R}{2}}{R} \right)^2 \right)^{-\frac{3}{2}} + \left(1 + \left(\frac{z - \frac{R}{2}}{R} \right)^2 \right)^{-\frac{3}{2}} \right) \vec{u}_z$$

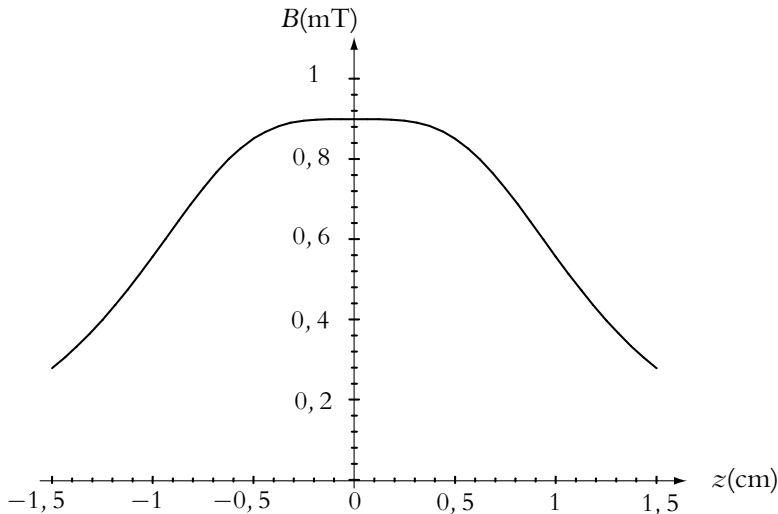


Figure 50.7 Norme du champ magnétique créé sur l'axe des bobines de Helmholtz pour $R = 1 \text{ cm}$ et $I = 10 \text{ A}$.

De plus :

$$\begin{aligned} \frac{B(0) - B\left(\pm \frac{d}{2}\right)}{B(0)} &= \frac{\frac{\mu_0 I}{R} \left(\frac{5}{4}\right)^{-\frac{3}{2}} - \frac{\mu_0 I}{2R} \left(1 + 2^{-\frac{3}{2}}\right)}{\frac{\mu_0 I}{R} \left(\frac{5}{4}\right)^{-\frac{3}{2}}} \\ &= \frac{2 \left(\frac{5}{4}\right)^{-\frac{3}{2}} - \left(1 + 2^{-\frac{3}{2}}\right)}{\left(\frac{5}{4}\right)^{-\frac{3}{2}}} = 5,42 \% \end{aligned}$$

Le champ entre les spires ne varie pas plus de 5 % par rapport à la valeur à mi-distance : on peut considérer que le champ est uniforme entre les deux spires.

4. Bobine torique

Une bobine torique est constituée de N spires identiques déduites les unes des autres par une rotation d'angle $\frac{2\pi}{N}$ autour de l'axe Oz .

La figure ci-contre donne l'exemple d'une bobine torique à section rectangulaire.

On n'a représenté que quelques spires pour assurer la clarté du schéma.

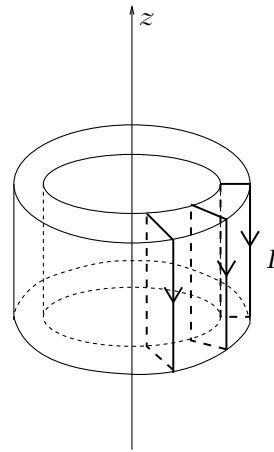


Figure 50.8 Bobine torique.

4.1 Étude des invariances

La distribution est invariante par toute rotation autour de l'axe Oz (on suppose que N est suffisamment grand pour cela ou que les spires sont régulièrement espacées).

On se place en coordonnées cylindriques.

$$||\vec{B}(M)|| = B(r, z)$$

4.2 Étude des symétries

Le plan passant par M et contenant l'axe Oz est un plan de symétrie de la distribution de courants : le champ magnétique est perpendiculaire à ce plan, soit :

$$\vec{B}(M) = B(r, z)\vec{u}_\theta$$

4.3 Calcul du champ par le théorème d'Ampère

On considère un cercle d'axe Oz de rayon r passant par M . La circulation du champ le long de ce contour est : $2\pi r B(r, z)$ et le courant enlacé vaut $I_{\text{enlacé}} = 0$ si M est à l'extérieur du tore (en M_1 ou M_3 par exemple) et $I_{\text{enlacé}} = NI$ si M est à l'intérieur (en M_2 par exemple). On en déduit :

$$2\pi r B(r, z) = \begin{cases} 0 & \text{pour } M \text{ à l'intérieur} \\ \mu_0 NI & \text{pour } M \text{ à l'extérieur} \end{cases}$$

et

$$\vec{B}(M) = \begin{cases} \vec{0} & \text{pour } M \text{ à l'intérieur} \\ \frac{\mu_0 NI}{2\pi r} \vec{u}_\theta & \text{pour } M \text{ à l'extérieur} \end{cases}$$

On remarque que le champ magnétique ne dépend pas de z ce qu'on ne pouvait pas prévoir *a priori*.

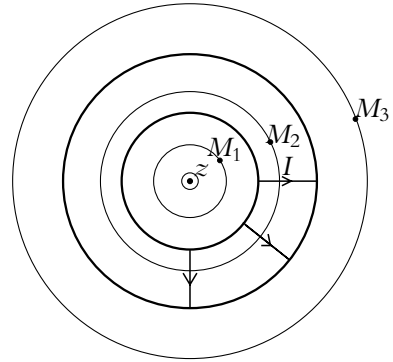


Figure 50.9 Contour d'Ampère pour une bobine torique.

5. Solénoïde

5.1 Calcul du champ magnétique sur l'axe

Un solénoïde est formé de N spires identiques, de même axe Oz , parcourues par un même courant, dans le même sens.

On définit le nombre de spires par unité de longueur : $n = \frac{N}{L}$ en notant L la longueur du solénoïde.

On choisit des spires circulaires.

On va calculer le champ magnétique créé sur l'axe du solénoïde.

Les propriétés d'invariance et de symétrie sont les mêmes que pour une spire unique : le champ est donc parallèle à l'axe et ne dépend que de la position du point sur l'axe : $\vec{B}(M) = B(z)\vec{u}_z$.

On utilise alors le résultat obtenu pour une spire et on applique le principe de superposition.

On note $z = OP$. La tranche d'épaisseur dz du solénoïde comprise entre z et $z + dz$ est parcourue par une intensité élémentaire $dI = nIdz$ et crée le champ élémentaire :

$$dB = n \frac{\mu_0 I}{2R} \sin^3 \alpha dz$$

avec $\tan \alpha = \frac{R}{PM}$.

$$PM = OM - OP = OM - z = \frac{R}{\tan \alpha} \quad \text{donc} \quad dz = \frac{R}{\sin^2 \alpha} d\alpha$$

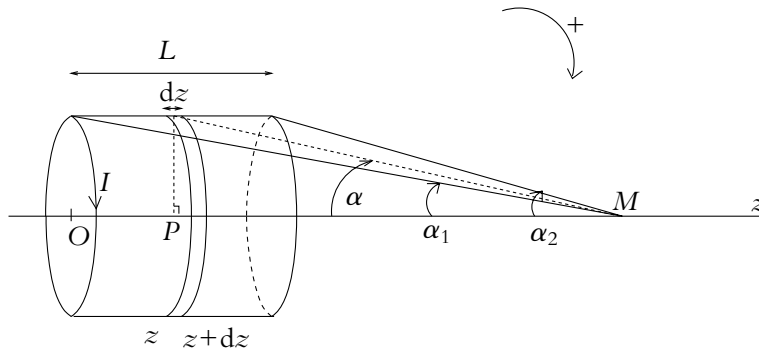


Figure 50.10 Solénoïde.

D'où

$$dB = n \frac{\mu_0 I}{2R} \sin^3 \alpha \frac{R}{\sin^2 \alpha} d\alpha = \frac{\mu_0 n I}{2} \sin \alpha d\alpha = -\frac{\mu_0 n I}{2} d(\cos \alpha)$$

On en déduit par intégration entre α_1 et α_2 :

$$B = \frac{\mu_0 n I}{2} (\cos \alpha_1 - \cos \alpha_2)$$

Avec l'origine de l'axe des z au milieu du solénoïde, on a :

$$\cos \alpha_1 = \frac{z + \frac{L}{2}}{\sqrt{R^2 + \left(z + \frac{L}{2}\right)^2}} \quad \text{et} \quad \cos \alpha_2 = \frac{z - \frac{L}{2}}{\sqrt{R^2 + \left(z - \frac{L}{2}\right)^2}}.$$

On en déduit :

$$B(z) = \frac{\mu_0 n I}{2} \left(\frac{z + \frac{L}{2}}{\sqrt{R^2 + \left(z + \frac{L}{2}\right)^2}} - \frac{z - \frac{L}{2}}{\sqrt{R^2 + \left(z - \frac{L}{2}\right)^2}} \right)$$

5.2 Analyse des variations le long de l'axe

Pour étudier les variations du champ le long de l'axe, on préfère une analyse graphique aux calculs (toujours possibles). L'allure de la norme du champ magnétique sur l'axe en fonction de z est la suivante :

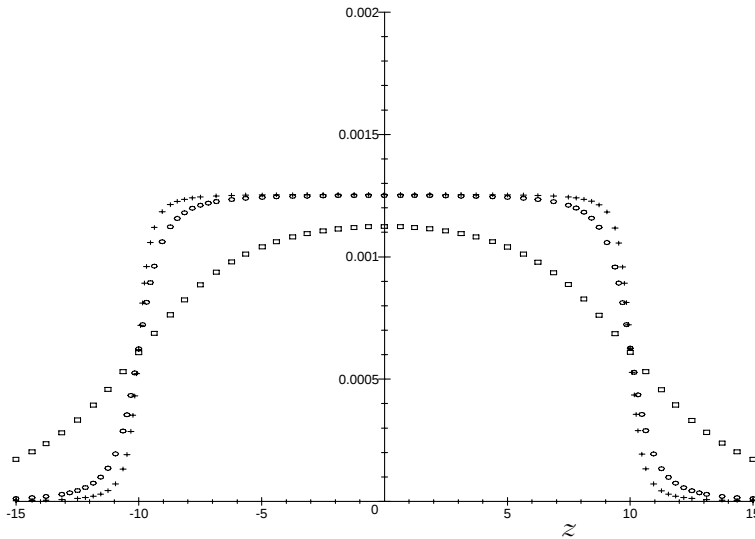


Figure 50.11 Module du champ magnétique le long de l'axe d'un solénoïde.

Sur la figure, on a pris : $L = 20$ cm et les valeurs $R = 5$ cm pour la courbe avec des carrés, $R = 1$ cm pour la courbe avec des losanges et $R = 0,5$ cm pour la courbe avec des croix.

Pour $R = 1$ cm ou $R = 0,5$ cm c'est-à-dire pour $\frac{R}{L} < 5 \cdot 10^{-2}$, le champ magnétique est quasi-uniforme à l'intérieur du solénoïde. Pour $R = 5$ cm c'est-à-dire pour $\frac{R}{L} = 0,25$, ce n'est plus vrai sur les bords du solénoïde.

Une autre approche pour savoir si les variations du champ sont importantes à l'intérieur du solénoïde consiste à comparer les valeurs en $z = 0$ (au centre du solénoïde) et $z = \frac{L}{2}$ (au bord du solénoïde) :

$$\left\{ \begin{array}{l} B(0) = \frac{\mu_0 n I}{2} \left(\frac{\frac{L}{2}}{\sqrt{R^2 + \frac{L^2}{4}}} + \frac{\frac{L}{2}}{\sqrt{R^2 + \frac{L^2}{4}}} \right) = \mu_0 n I \frac{L}{\sqrt{4R^2 + L^2}} \\ B\left(\pm \frac{L}{2}\right) = \frac{\mu_0 n I}{2} \left(0 + \frac{L}{\sqrt{R^2 + L^2}} \right) = \mu_0 n I \frac{L}{\sqrt{4R^2 + 4L^2}} \end{array} \right.$$

Donc

$$\frac{B\left(\pm\frac{L}{2}\right)}{B(0)} = \frac{1}{2} \sqrt{\frac{4R^2 + L^2}{R^2 + L^2}}$$

Ce rapport prend les valeurs suivantes en fonction du rapport entre R et L :

$\frac{L}{R}$	0,5	1	2	5	10	20	100
$\frac{B(\pm\frac{L}{2})}{B(0)}$	0,92	0,79	0,63	0,53	0,51	0,50	0,50

L'allure de $\frac{B\left(\pm\frac{L}{2}\right)}{B(0)}$ en fonction de $x = \frac{L}{R}$ est représentée sur la figure 50.12.

L'amplitude du champ sur l'axe ne diminue jamais de plus de la moitié entre le centre et les extrémités du solénoïde car

$$\lim_{\frac{L}{R} \rightarrow \infty} \frac{B\left(\pm\frac{L}{2}\right)}{B(0)} = \frac{1}{2}$$

On pourra donc considérer le champ à l'intérieur du solénoïde comme uniforme lorsqu'on n'est pas trop près des bords.

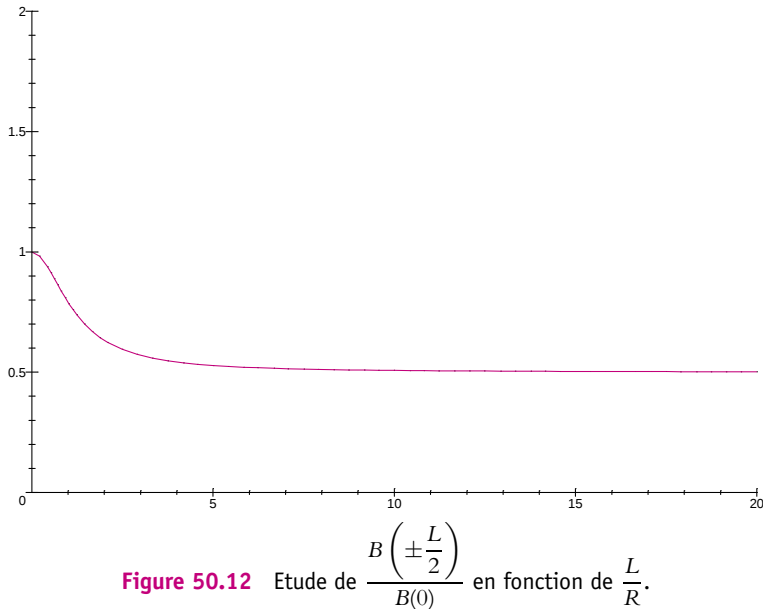


Figure 50.12 Etude de $\frac{B\left(\pm\frac{L}{2}\right)}{B(0)}$ en fonction de $\frac{L}{R}$.

5.3 Topographie du solénoïde de longueur finie

La topographie du champ magnétique d'un solénoïde de longueur finie est la suivante :

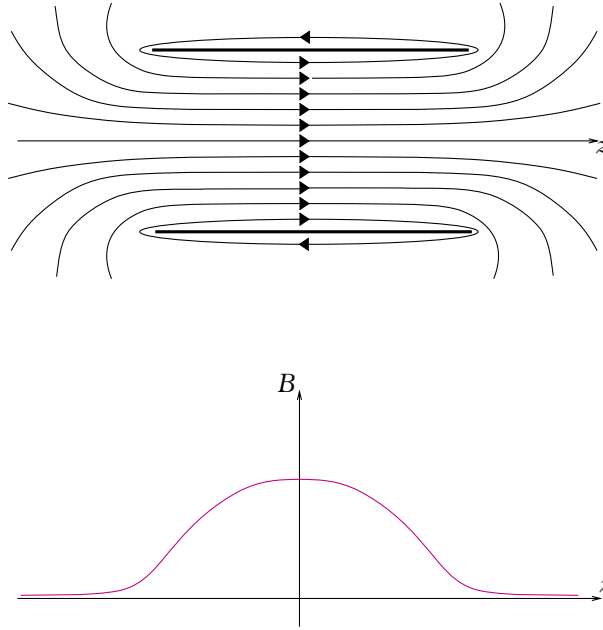


Figure 50.13 Topographie du champ magnétique créé par un solénoïde.

5.4 Cas d'un solénoïde infini

Soit un solénoïde de longueur infinie, de section circulaire.

a) Champ sur l'axe

Pour obtenir le champ sur l'axe, il suffit de prendre $\alpha_1 = 0$ et $\alpha_2 = \pi$ dans l'expression obtenue pour un solénoïde de longueur finie. On obtient :

$$B = \frac{\mu_0 n I}{2} (1 + 1) = \mu_0 n I$$

soit

$$\vec{B} = \mu_0 n I \vec{u}_z$$

Le champ est alors indépendant du rayon du solénoïde.

b) Champ en tout point

Le système est invariant par translation le long de l'axe Oz du solénoïde : on se place en coordonnées cylindriques et on a :

$$||\vec{B}(M)|| = B(r, \theta)$$

La section du solénoïde est circulaire, le système est donc invariant par rotation autour de l'axe Oz , on a :

$$||\vec{B}(M)|| = B(r)$$

Le plan perpendiculaire à l'axe Oz et passant par M est un plan de symétrie : le champ magnétique est perpendiculaire à ce plan soit

$$\vec{B}(M) = B(r)\vec{u}_z$$

On peut alors appliquer le théorème d'Ampère pour établir que le champ ne prend que deux valeurs possibles suivant que le point M est à l'intérieur ou à l'extérieur du solénoïde et pour obtenir les valeurs correspondantes du champ.

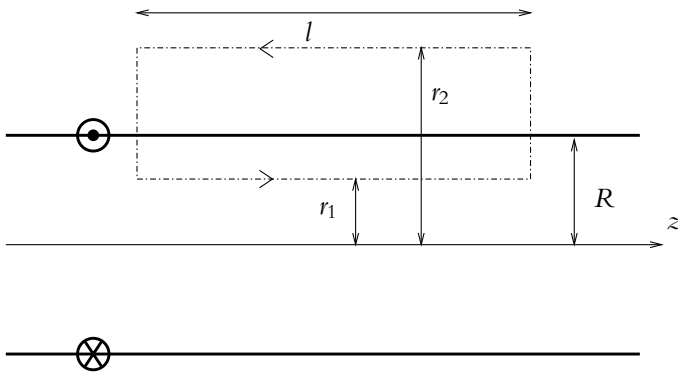


Figure 50.14 Contour d'Ampère sur un solénoïde infini.

Soit un contour $\mathcal{C}$ fermé contenu dans un plan contenant l'axe Oz et constitué de deux parallèles à Oz distantes respectivement de r_1 et r_2 de l'axe et de deux perpendiculaires à Oz .

Les contributions des deux perpendiculaires à Oz à la circulation du champ magnétique le long de ce contour sont nulles : le champ est perpendiculaire aux déplacements élémentaires.

En notant alors l la longueur des parallèles à Oz , on obtient :

$$\oint_{M \in \mathcal{C}} \vec{B}(M) \cdot d\vec{OM} = B(r_1)l - B(r_2)l = \mu_0 I_{\text{enlacé}}$$

On distingue les trois cas suivants :

- Si $r_1 < r_2 < R$, $I_{\text{enlacé}} = 0$ et on en déduit :

$$B(r_1) = B(r_2)$$

le champ magnétique créé à l'intérieur d'un solénoïde infini ne dépend pas de r : il est donc uniforme. On le note B_{int} .

- Si $R < r_1 < r_2$, $I_{\text{enlacé}} = 0$ et on en déduit :

$$B(r_1) = B(r_2)$$

le champ magnétique créé à l'extérieur d'un solénoïde infini est également uniforme. On le note B_{ext} .

- Si $r_1 < R < r_2$ (cf. figure 50.14), $I_{\text{enlacé}} = nI$ et le champ magnétique est orienté dans le sens de $\vec{u}_z$. Compte tenu de l'orientation du contour donné sur la figure 50.14, on en déduit :

$$B(r_1)l - B(r_2)l = (B_{\text{int}} - B_{\text{ext}})l = \mu_0 nI$$

Donc

$$B_{\text{int}} - B_{\text{ext}} = \mu_0 nI$$

Dans le cas d'un solénoïde de section circulaire, on a calculé la valeur du champ sur l'axe et on peut écrire :

$$B_{\text{int}} = \mu_0 nI$$

et en déduire :

$$B_{\text{ext}} = 0$$

Si on n'a pas au préalable calculé le champ sur l'axe, il suffit de considérer que le solénoïde est une bobine torique de rayon infini. En utilisant les résultats obtenus pour une bobine torique, on en déduit que

$$B_{\text{ext}} = 0$$

et par la relation liant B_{ext} et B_{int} que

$$B_{\text{int}} = \mu_0 nI$$

On remarque que dans le cas d'une section de forme quelconque, le champ magnétique vérifie :

$$\vec{B} = B(r, \theta) \vec{u}_z$$

En choisissant des contours dans un plan $\theta = \text{constante}$, le calcul précédent à partir de la bobine torique est valable. Le résultat est donc indépendant de la forme de la section du solénoïde.

6. Complément : nappe plane de courant

Soit une distribution surfacique de courants uniforme $\vec{j}_s = j_s \vec{u}_y$ circulant dans le plan $x = 0$ ou encore yOz .

Cette distribution peut être considérée comme la superposition d'une infinité de fils infinis parallèles les uns aux autres.

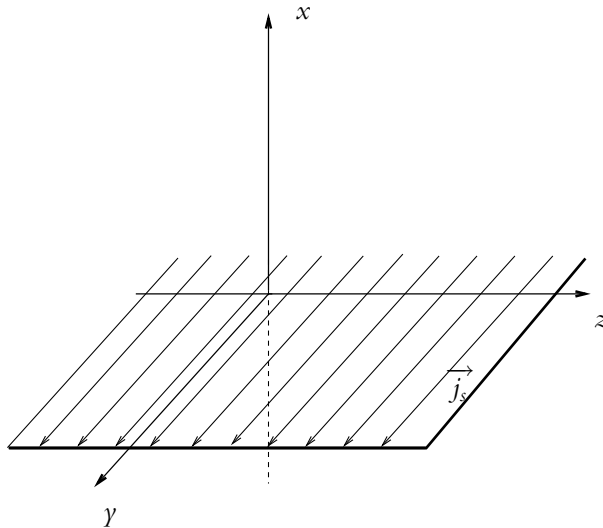


Figure 50.15 Nappe plane de courants.

6.1 Étude des invariances

La distribution étant infinie, elle est invariante par translation selon Oy et selon Oz , donc le champ magnétique ne dépend, en coordonnées cartésiennes, que de la distance x du point au plan $x = 0$.

6.2 Étude des symétries

Le plan passant par M et dirigé par $\vec{u}_x$ et $\vec{u}_y$ est un plan de symétrie de la distribution donc le champ magnétique lui est perpendiculaire.

On en déduit finalement que :

$$\vec{B}(M) = B(x) \vec{u}_z$$

Le plan de la distribution yOz est plan de symétrie donc les propriétés de symétrie imposent les orientations suivantes pour le champ en M et en M' , symétrique de M par rapport au plan $x = 0$:

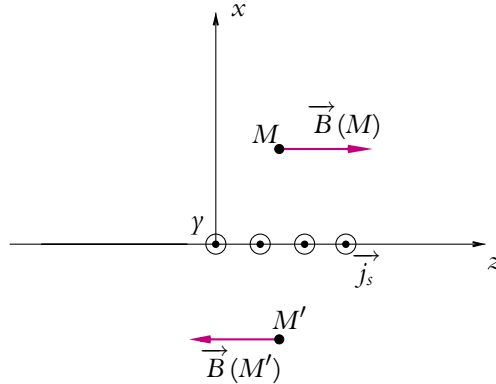


Figure 50.16 Nappe plane de courants et symétries par rapport à un plan.

On en déduit :

$$B(x) = -B(-x)$$

6.3 Calcul du champ par le théorème d'Ampère

On va calculer la norme du champ en appliquant le théorème d'Ampère à un contour fermé constitué de deux parallèles à Ox et de deux parallèles à Oz situées en x et en $-x$. On note l la longueur du contour parallèle à Oz .

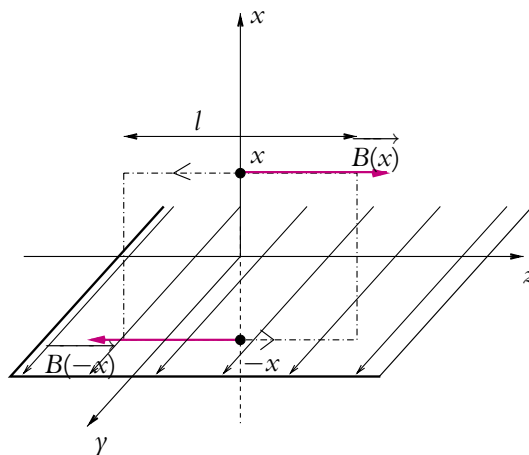


Figure 50.17 Contour d'Ampère d'une nappe plane de courants.

Les contributions sur les parallèles à Ox à la circulation du champ sont nulles car le champ est perpendiculaire à Ox .

Le théorème d'Ampère s'écrit donc :

$$-B(x)l + B(-x)l = \mu_0 I_{\text{enlacé}}$$

Le courant enlacé par le contour est : $I_{\text{enlacé}} = j_s l$. On notera que le choix de l'orientation du contour détermine le sens du vecteur surface associé donc le signe de $I_{\text{enlacé}}$. Un changement d'orientation change à la fois le signe de la circulation et celui de $I_{\text{enlacé}}$.

On en déduit :

$$-B(x)l + B(-x)l = -2B(x)l = I_{\text{enlacé}} = j_l$$

donc :

$$\overrightarrow{B(M)} = -\frac{\mu_0 j_s}{2} \overrightarrow{u_z} = \frac{\mu_0}{2} \overrightarrow{j_s} \wedge \overrightarrow{u_x}$$

$\overrightarrow{u_x}$ est le vecteur directeur de la normale à la surface. Cette dernière formulation est intéressante par son caractère intrinsèque et donc son indépendance vis-à-vis du repère cartésien initialement choisi.

Dans ce cas, les lignes de champ sont donc des parallèles à Oz parcourues dans le sens décroissant pour $x > 0$ et croissant pour $x < 0$.

On note que le champ est uniforme de part et d'autre de la nappe, ce qui était déjà annoncé par le théorème d'Ampère le long d'un contour bien choisi (c'est-à-dire un rectangle ne coupant pas le plan $x = 0$).

D'autre part, en notant $\overrightarrow{B}_+$ et $\overrightarrow{B}_-$ les valeurs du champ respectivement pour $x > 0$ et $x < 0$, on obtient :

$$\overrightarrow{B}_+ - \overrightarrow{B}_- = \mu_0 \overrightarrow{j_s} \wedge \overrightarrow{u_x}$$

On généralisera ce résultat dans le cours de seconde année.

Le programme officiel de première année limite l'étude aux champs magnétiques créés par des circuits filiformes. Les exercices portant sur des distributions surfaciques et volumiques sont proposés à titre de complément et d'approfondissement.

A. Applications directes du cours

1. Champ magnétique créé par une nappe volumique

On considère la distribution infinie suivant Ox et Oy :

$$\begin{cases} \vec{j} = j\vec{u}_x & \text{pour } |z| < a \\ \vec{j} = \vec{0} & \text{pour } |z| > a \end{cases}$$

1. Étudier les invariances de ce système.
2. Analyser ses propriétés de symétrie.
3. En déduire qu'on peut appliquer le théorème d'Ampère pour déterminer le champ magnétique.
4. Déterminer le champ magnétique en tout point de l'espace.

2. Solénoïde épais

Un solénoïde épais infiniment long d'axe Oz , est parcouru par des courants orthoradiaux de densité volumique uniforme $\vec{j} = j\vec{u}_\theta$ pour $a < r < b$. Les coordonnées utilisées sont les coordonnées cylindro-polaires (ρ, θ, z) de vecteur $(\vec{u}_\rho, \vec{u}_\theta, \vec{u}_z)$.

Calculer $\vec{B}$ en tout point de l'espace.

3. Champ tournant

1. Soient deux solénoïdes de n spires par unité de longueur, de longueur L et de rayon R distants de $2l$, de même axe et parcourus par le même courant I dans le même sens. Calculer le champ magnétique B au centre du système. On pose $B = kI$. Expliciter k .

Application numérique : $R = 3 \text{ cm}$, $L = 5 \text{ cm}$, $n = 5000$, $l = 7,5 \text{ cm}$.

2. On dispose de trois paires de tels solénoïdes orientés à 120° et alimentés en série selon le schéma suivant :

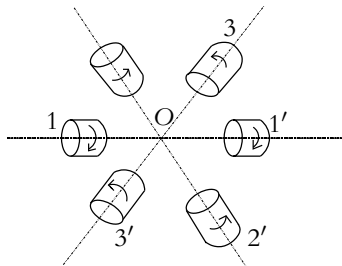


Figure 50.18

Calculer le champ magnétique au centre et tracer l'allure des lignes de champ.

3. Les trois paires de solénoïdes sont alimentées par trois courants sinusoïdaux déphasés entre eux de $\frac{2\pi}{3}$:

$$i_1 = I\sqrt{2} \cos \omega t, \quad i_2 = I\sqrt{2} \cos \left(\omega t - \frac{2\pi}{3} \right), \quad i_3 = I\sqrt{2} \cos \left(\omega t - \frac{4\pi}{3} \right)$$

En admettant que l'expression du champ magnétique reste inchangée avec un courant d'intensité variable, montrer qu'on obtient au centre un champ magnétique tournant dont on calculera le module en fonction de k et I . Quelle est la vitesse de rotation du vecteur $\vec{B}$?

B. Exercices et problèmes

1. Distributions cylindriques, d'après ENAC 2001

Un cylindre de révolution autour de l'axe Oz a pour rayon b et une longueur « infinie » (très grande devant b). Il est parcouru dans la direction et dans le sens de Oz par un courant continu de densité uniforme de courant $\vec{J}$.

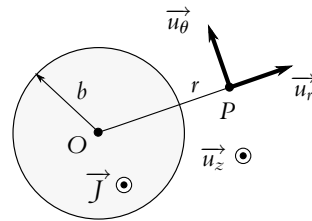


Figure 50.19

1. a) Par analyse des invariances de la distribution de courant, déterminer la dépendance du champ $\vec{B}$ en coordonnées cylindriques (r, θ, z) , pour un point P quelconque de l'espace.

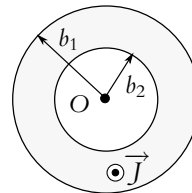
b) Par analyse des symétries de la distribution de courant, déterminer la direction du champ $\vec{B}$ dans la base cylindrique $(\vec{u}_r, \vec{u}_\theta, \vec{u}_z)$ pour un point P quelconque de l'espace.

2. a) Déterminer le vecteur champ magnétique $\vec{B}$ créé par ce courant en un point P extérieur au cylindre, situé à la distance r de Oz .

b) Même question lorsque P est à l'intérieur du cylindre.

c) Donner une expression vectorielle intrinsèque (indépendante du système de coordonnées) du vecteur champ calculé dans la question précédente.

3. Un cylindre « de longueur infinie » et de révolution autour de l'axe Oz est creux ; la partie pleine est comprise entre les rayons b_1 et b_2 ($b_1 > b_2$). Elle est parcourue dans la direction et dans le sens de Oz par un courant continu de densité uniforme de courant $\vec{J}$.



a) Déterminer le vecteur champ $\vec{B}$ en un point P tel $r(P) < b_2$. Figure 50.20

b) Déterminer le vecteur champ $\vec{B}$ en un point P tel $b_2 < r(P) < b_1$.

4. Un cylindre de longueur « infini » et de révolution autour de l'axe O_1z a pour rayon b_1 . On creuse dans le cylindre un autre cylindre de "longueur infinie" et de révolution autour de l'axe O_2z parallèle à O_1z et de même sens ; son rayon est $b_2 < b_1$.

On désigne par $2a$ la distance O_1O_2 . Dans la partie pleine cercle dans la direction et le sens de O_1z un courant continu de densité uniforme J .

En appliquant le principe de superposition et en utilisant l'expression du champ créé par un cylindre, déterminer l'expression du champ magnétique à l'intérieur de la cavité.

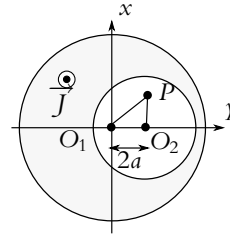


Figure 50.21

2. Bobine torique, d'après ENAC 2007

Parmi les quatre propositions pour chaque question, il peut y avoir une ou plusieurs réponses justes.

Une bobine est constituée par un fil conducteur bobiné en spires jointives sur un tore circulaire à section carrée de côté a de rayon moyen R . On désigne par n le nombre total de spires et par I le courant qui les parcourt.

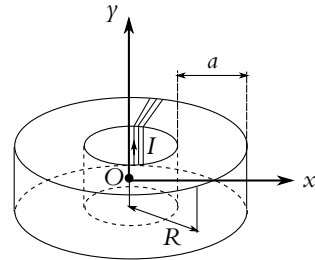


Figure 50.22

1. Tout plan méridien du bobinage c'est-à-dire tout plan contenant l'axe de révolution Oy est :

- A) Plan de symétrie de la distribution de courant.
- B) Plan d'antisymétrie de la distribution de courant.
- C) Plan d'antisymétrie du champ magnétique.
- D) Plan de symétrie du champ magnétique.

2. Il en résulte que les lignes de champ du champ magnétique passant par un point quelconque M situé à l'intérieur de la bobine sont :

- A) Des cercles d'axe Oy .
- B) Des cercles de centre O .
- C) Des cercles dont le centre est le centre de la spire contenant M .
- D) Des carrés dont l'un des sommets contient M .

3. Calculer la norme du champ magnétique qui règne en un point $M(x, y)$ quelconque du plan xOy à l'intérieur du tore.

$$A) B = \frac{\mu_0 n I}{2\pi \sqrt{x^2 + y^2}} \quad B) B = \frac{\mu_0 n I}{2\pi x}$$

$$C) B = \frac{\mu_0 n I}{2\pi y} \quad D) B = \frac{\mu_0 n I x}{2\pi \sqrt{x^2 + y^2}}$$

4. Calculer le flux Φ du champ magnétique à travers la surface d'une spire dont la normale est orientée dans le sens du champ.

$$A) \Phi = \frac{\mu_0 n I a}{2\pi} \ln \frac{2R + a}{2R - a} \quad B) \Phi = \frac{\mu_0 n I a}{2\pi} \ln \frac{R + a}{R - a}$$

$$C) \Phi = \frac{\mu_0 n I a}{2\pi} \ln \frac{R + 2a}{2R - 2a} \quad D) \Phi = \frac{\mu_0 n I a}{2\pi} \ln \frac{R}{a}$$

5. On désigne respectivement par $B_{\min}$ et $B_{\max}$ les valeurs minimum et maximum du champ magnétique à l'intérieur de la bobine. Calculer la valeur numérique du rapport a/R pour une variation relative du champ de 10% :

$$A) \frac{a}{R} = 0,100 \quad B) \frac{a}{R} = 0,050$$

$$C) \frac{a}{R} = 0,075 \quad D) \frac{a}{R} = 0,200$$

3. Rubans infinis

On considère un fil infini, sans épaisseur, parcouru par un courant I . On considère que le fil constitue l'axe Ox , l'intensité I étant dans le sens des x croissants.

1. Déterminer le champ magnétique $\vec{B}$ créé par le fil en tout point de l'espace.

2. On s'intéresse maintenant à un ruban infini selon Ox , d'épaisseur nulle selon Oz et de largeur $2a$ ($-a \leq y \leq a$) selon Oy (figure 50.23). Ce ruban est parcouru par un intensité I_1 dans le sens des x croissants.

a) En décomposant le ruban en fils parallèles de largeur dy , exprimer l'intensité dI , parcourant l'un de ces fils élémentaires en fonction de I_1 , dy et a .

b) On souhaite calculer le champ magnétique créé par le ruban en un point M de l'axe Oz .

1. Par analyse des symétries de la distribution de courant, déterminer la direction du champ $\vec{B}$ en M .

2. En utilisant l'expression de $\vec{B}$ créé par un fil infini, déterminer l'expression du champ créé en M par le ruban, en fonction de a , z cote de M , μ_0 et I_1 .

c) On dispose un deuxième ruban, identique et parallèle au premier, à une distance b au-dessus. L'intensité parcourant ce second ruban est aussi I_1 mais selon les x décroissants. Calculer le champ résultant en un point de l'axe Oz situé à mi-chemin entre les rubans.

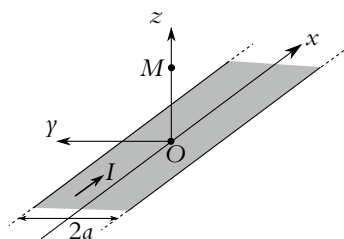


Figure 50.23

4. Disque de Rowland

Un disque circulaire D , de centre O , de rayon R , porte la charge surfacique constante σ . On fait tourner D à la vitesse angulaire ω constante autour de Oz . Les charges sont entraînées

dans le mouvement et fixes par rapport au disque (la densité surfacique de charge σ n'est pas modifiée).

1. Déterminer la densité surfacique de courant équivalente au disque en rotation.
2. On décompose le disque en spires élémentaires de rayon r , d'épaisseur dr :
 - a) Calculer le courant dI dans une spire.
 - b) Calculer le champ $\vec{B}$ en tout point de l'axe Oz .

5. Enroulement tronconique, d'après ENAC 2000

On réalise un bobinage en enroulant sur un tronc de cône, jointivement suivant la génératrice, N spires d'un fil de cuivre de diamètre a et de résistivité ρ (figure 50.24). Le tronc de cône de sommet S , de demi-angle au sommet α , est caractérisé par les rayons r_1 et $r_2 > r_1$ de ses deux bases.

Chaque spire est repérée par sa cote z qui mesure la distance qui sépare son centre de S . On désigne par r le rayon de la spire située à la cote z . Le vecteur unitaire de l'axe Sz est $\vec{u}_z$.

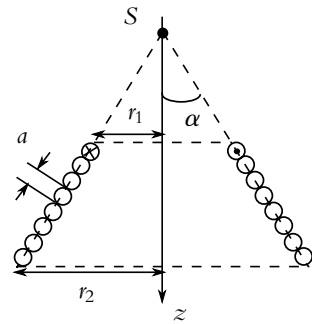


Figure 50.24

1. Exprimer le nombre N de spires qui constituent le bobinage en fonction de r_1 , r_2 , a et α .

$$A) N = \frac{r_2 - r_1}{a \cos \alpha} \quad B) N = \frac{r_2 - r_1}{a \tan \alpha}$$

$$C) N = \frac{r_2 + r_1}{2a \cos \alpha} \quad D) N = \frac{r_2 - r_1}{a \sin \alpha}$$

2. On désigne par dN le nombre de spires dont la cote est comprise entre z et $z + dz$. On considère que ces dN spires ont la même circonférence et qu'elles créent le même champ magnétique. Exprimer dN .

$$A) dN = \frac{dz}{a \cos \alpha} \quad B) dN = \frac{dz}{a \sin \alpha}$$

$$C) dN = \frac{dz}{a \tan \alpha} \quad D) dN = \frac{dz}{2a \sin \alpha}$$

3. La résistance R d'un fil de résistivité ρ , de section s et de longueur ℓ est donnée par la relation : $R = \rho \ell / s$. Calculer R .

$$A) R = \rho \frac{r_2^2 - r_1^2}{a^3 \cos \alpha} \quad B) R = 4\rho \frac{r_2^2 - r_1^2}{a^3 \sin \alpha}$$

$$C) R = 2\rho \frac{r_2^2 - r_1^2}{a^3 \tan \alpha} \quad D) R = \rho \frac{r_2^2 + r_1^2}{2a^3 \cos \alpha}$$

4. Le bobinage est parcouru par un courant I dans le sens représenté sur la figure (50.24). On désigne par μ_0 la perméabilité du vide. Calculer le champ magnétique $\vec{B}_1$ créé en S par une spire de rayon r .

$$\begin{aligned} A) \vec{B}_1 &= \frac{\mu_0 I}{2r} \sin^3 \alpha \vec{u}_z & B) \vec{B}_1 &= \frac{\mu_0 I}{r} \sin^3 \alpha \vec{u}_z \\ C) \vec{B}_1 &= \frac{\mu_0 I}{2\pi r} \sin^3 \alpha \vec{u}_z & D) \vec{B}_1 &= \frac{\mu_0 I}{4\pi r} \sin^3 \alpha \vec{u}_z \end{aligned}$$

5. En déduire le champ magnétique créé en S par la totalité du bobinage.

$$\begin{aligned} A) \vec{B} &= \frac{\mu_0 I \sin^3 \alpha}{2\pi a} \ln \frac{r_2 - r_1}{r_2 + r_1} \vec{u}_z & B) \vec{B} &= \frac{\mu_0 I \sin^3 \alpha}{2\pi(r_2 - r_1)} \ln \frac{r_2}{r_1} \vec{u}_z \\ C) \vec{B} &= \frac{\mu_0 I \sin^2 \alpha}{2a} \ln \frac{r_2}{r_1} \vec{u}_z & D) \vec{B} &= \frac{\mu_0 I \sin^2 \alpha}{2\pi a} \ln \frac{r_2 - r_1}{r_2 + r_1} \vec{u}_z \end{aligned}$$

6. Câble coaxial, d'après C.C.P. TSI 2002

1. Énoncer le théorème d'Ampère relatif à la circulation du champ magnétique $\vec{B}$ le long d'un contour fermé $\mathcal{C}$ constitué de points M et s'appuyant sur une surface S . On notera $\vec{j}(P)$ la densité de courant en un point P de la surface S .

On considère un câble coaxial cylindrique de longueur supposée infinie, constitué d'un conducteur central plein de rayon R_1 , parcouru par un courant uniforme d'intensité I et d'un conducteur périphérique évidé, de rayon intérieur R_2 , de rayon extérieur R_3 avec $R_1 < R_2 < R_3$ et parcouru par un courant uniforme également d'intensité I mais circulant en sens inverse par rapport au courant conducteur central.

On notera $\vec{u}_z$ le vecteur directeur unitaire de l'axe commun des 2 conducteurs. Soit un point M à une distance r de l'axe du câble.

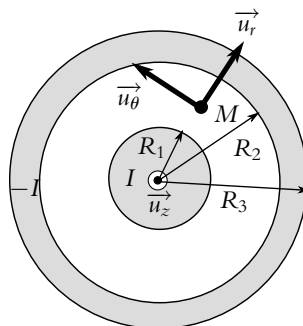


Figure 50.25

2. a) Montrer que le champ magnétique $\vec{B}$ créé au point M est orthoradial.
- b) Montrer qu'il peut se mettre sous la forme $\vec{B} = B(r)\vec{u}_\theta$ où $B(r)$ est une fonction de r uniquement.
- c) Préciser alors la forme des lignes de champ.
- 3.
- a) Montrer que le champ magnétique $\vec{B}$ créé au point M est nul si $r > R_3$
- b) Expliquer l'intérêt du câble coaxial par rapport à un fil simple parcouru par un courant de même intensité.
4. Calculer les densités de courant $\vec{j}_1$ et $\vec{j}_2$ respectivement du conducteur central et du conducteur périphérique en fonction des courants I_1 et I_2 et des rayons R_1 , R_2 et R_3 .
5. En appliquant le théorème d'Ampère à un contour $\mathcal{C}$ que l'on précisera, donner l'expression de la composante $B(r)$ du champ magnétique créé au point M en fonction de μ_0 , I , r , R_1 , R_2 et R_3 dans chacun des trois cas suivants :
- a) $r < R_1$
- b) $R_1 < r < R_2$
- c) $R_2 < r < R_3$
6. Vérifier que le champ magnétique est continu en $r = R_1$ et en $r = R_2$.
- a) Tracer le graphe de $B(r)$.

51

Dipôle magnétique (PCSI)

Ce chapitre envisage un analogue du dipôle électrostatique : le dipôle magnétique pour lequel on obtient des résultats en tout point similaires à ceux de l'électrostatique, que ce soit pour le champ créé ou l'action subie dans un champ extérieur.

Un des intérêts de l'étude de ces dipôles magnétiques vient de l'interprétation microscopique des aimants : on assimile les sources microscopiques de champ magnétique à des dipôles magnétiques.

1. Définitions

1.1 Moment magnétique

Soit une boucle circulaire de courant d'intensité I .

On peut lui associer un vecteur surface orienté $\vec{S}$ dont le module correspond à la surface intérieure au cercle et dont la direction et le sens sont donnés par la règle du tire-bouchon en fonction du sens du courant dans la boucle.

On appelle alors *moment magnétique* de la spire la quantité :

$$\vec{\mathcal{M}} = I \vec{S} = IS \vec{n}$$

D'après cette définition, l'unité du moment magnétique est l'ampère mètre carré ($A.m^2$) dans les unités du Système International.

On peut généraliser cette définition à une boucle de courant de forme quelconque.

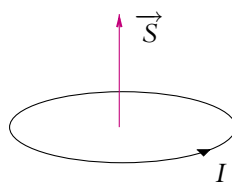


Figure 51.1 Définition du moment magnétique par une spire parcourue par un courant I .

1.2 Expérience d'Oersted

En 1820, Oersted observe qu'un aimant (une boussole par exemple) est dévié lorsqu'on approche un fil parcouru par un courant continu. D'autre part, l'aimant s'oriente par rapport au fil de manière orthoradiale.

Si on refait la même expérience en remplaçant l'aimant par une spire, les résultats sont identiques.

Par conséquent, un aimant se comporte d'un point de vue magnétique comme une spire parcourue par un courant. On pourra donc modéliser un aimant par une telle spire et une spire par un aimant. Cela conduit à l'idée d'interpréter les propriétés magnétiques d'un aimant en terme de petites boucles de courant.

1.3 Quelques ordres de grandeur de moment magnétique

Suivant l'échelle à laquelle on se place, une large plage d'ordres de grandeur est possible pour les moments magnétiques.

- Au niveau atomique ou nucléaire, les moments magnétiques sont quantifiés par le magnéton de Bohr ou par le magnéton nucléaire qui valent :
 - magnéton de Bohr :

$$\mu_B = \frac{e\hbar}{2m_e} = \frac{1,6 \cdot 10^{-19} 10^{-34}}{2 \times 9,1 \cdot 10^{-31}} = 9,27 \cdot 10^{-24} \text{ A.m}^2$$

correspondant au moment magnétique de l'électron,

- magnéton nucléaire :

$$\mu_N = \frac{e\hbar}{2m_p} = \frac{1,6 \cdot 10^{-19} 10^{-34}}{2 \times 1,6 \cdot 10^{-27}} = 5,05 \cdot 10^{-27} \text{ A.m}^2$$

correspondant au moment magnétique du proton.

- Pour une boussole, le moment magnétique est de l'ordre de 10 A.m^2 .
- Pour la Terre qu'on assimile à un dipôle magnétique, le moment magnétique vaut $7,5 \cdot 10^{22} \text{ A.m}^2$.

À propos de la boussole et de la Terre, on peut préciser que le pôle Nord de la Terre correspond au pôle Sud du dipôle magnétique qui modélise le comportement magnétique de la planète. En effet, on a historiquement appelé pôle Nord le pôle vers lequel se dirige le pôle Nord de la boussole. Ce dernier étant attiré par un pôle Sud, le pôle Nord de la Terre est un pôle Sud pour l'aimant correspondant.

1.4 Dipôle magnétique

On appelle *dipôle magnétique* toute distribution de courants permanents dont le moment magnétique est non nul et dont les dimensions sont faibles par rapport aux distances de la distribution aux autres éléments.

On adoptera comme modèle d'un dipôle magnétique une boucle de courant de surface orientée $\vec{S}$ et parcourue par un courant I : cette boucle présente un moment magnétique :

$$\vec{M} = I \vec{S}$$

On ne se préoccupera pas forcément de savoir s'il s'agit d'une distribution volumique ou d'une simple « petite » spire.

On s'intéresse dans la suite aux deux aspects suivants :

- au rôle actif du dipôle magnétique : le dipôle crée un champ magnétique,
- en complément, au rôle passif du dipôle magnétique : le dipôle subit l'action d'un champ magnétique extérieur.

On garde en mémoire que lors de l'interaction de deux dipôles ces derniers auront à la fois un rôle passif et un rôle actif.

2. Champ magnétique créé par un dipôle magnétique

On commence par étudier le rôle actif du dipôle en déterminant le champ créé par un dipôle magnétique modélisé par une boucle de courant dont on étudie le champ à grande distance. On suppose dans un premier temps que la boucle de courant est circulaire.

2.1 Champ magnétique créé par une boucle de courant

On a établi au paragraphe 3.3 du chapitre 50 que le champ magnétique créé par une spire circulaire le long de son axe s'écrit :

$$\vec{B} = \frac{\mu_0 I}{2R} \sin^3 \alpha \vec{u}_z$$

Pour établir ce résultat, on a utilisé le fait que la distance du point M où on cherche le champ à n'importe quel point de la distribution est constante.

En revanche, si on se place « loin » de la spire, on pourra faire un développement limité et obtenir des résultats approchés pour tout point « loin » de la spire. On retrouve les mêmes idées que lors de l'étude du dipôle électrostatique : faire un développement limité pour avoir une expression approchée « loin » de la distribution.

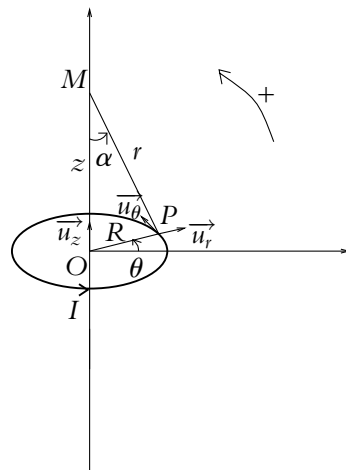


Figure 51.2 Boucle de courant.

2.2 Topographie du champ « loin » du dipôle magnétique

On peut déterminer expérimentalement la topographie du champ magnétique loin du dipôle (par exemple avec un aimant et de la limaille de fer permettant de visualiser les lignes de champ).

Expérimentalement on observe que, loin du dipôle, ces lignes de champ sont les mêmes que les lignes de champ électrostatique d'un dipôle électrostatique. Or le champ créé en M par un dipôle de moment $\vec{p}$ placé en O s'exprimait par :

$$\vec{E}(M) = \frac{1}{4\pi\epsilon_0} \frac{3(\overrightarrow{OM} \cdot \vec{p})\overrightarrow{OM} - OM^2\vec{p}}{OM^5}$$

Par analogie, on a donc formellement la même expression en remplaçant le moment dipolaire $\vec{p}$ par le moment magnétique $\vec{M}$ et la constante $\frac{1}{\epsilon_0}$ par μ_0 . La figure suivante donne l'allure des lignes de champ loin de la spire. Elle est invariante par rotation autour de l'axe du dipôle

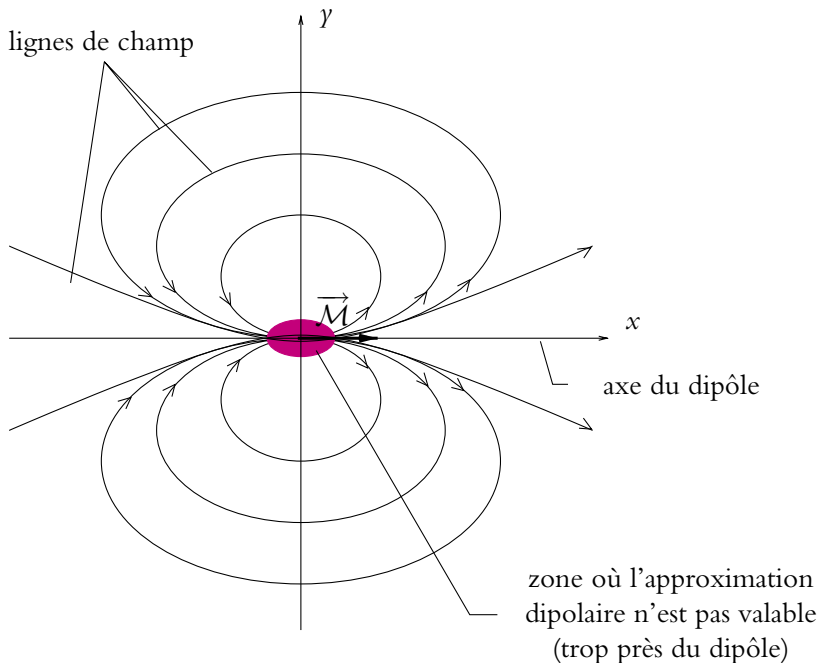


Figure 51.3 Topographie du champ magnétique créé par un dipôle magnétique.

L'analogie ne s'étend pas à ce qui se passe à proximité du dipôle.

En effet, le modèle du dipôle électrique est celui de deux charges opposées : les lignes de champ au voisinage du dipôle relient ces deux charges.

Quant au dipôle magnétique, on peut le modéliser par une spire : les lignes de champ « tournent » autour de la spire.

Au voisinage des dipôles, on obtient donc les cartes de champ suivantes :

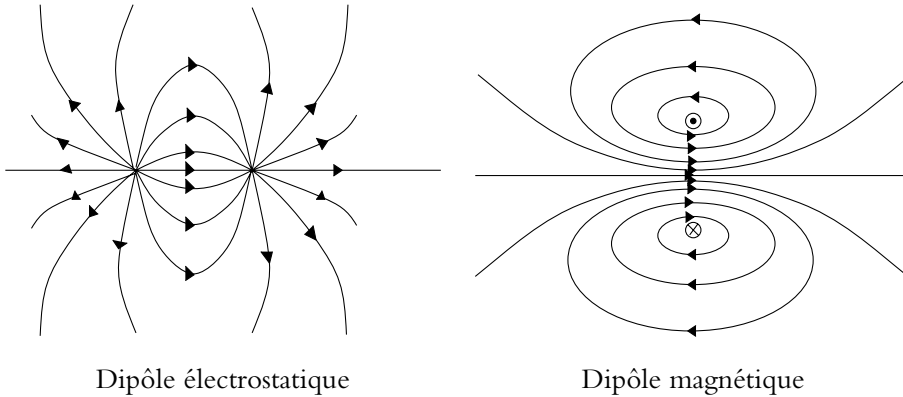


Figure 51.4 Comparaison des lignes de champ des dipôles électrique et magnétique au voisinage des dipôles.

2.3 Champ magnétique créé par un dipôle magnétique

On note N dans ce paragraphe le point où on détermine le champ créé par le dipôle magnétique. D'après le paragraphe précédent, on a :

$$\vec{B}(N) = \frac{\mu_0}{4\pi} \frac{3(\vec{ON} \cdot \vec{\mathcal{M}})\vec{ON} - ON^2 \vec{\mathcal{M}}}{ON^5}$$

ou encore, dans la base sphérique, l'axe Oz étant celui du dipôle :

$$\vec{B} = \frac{\mu_0}{4\pi r^3} \mathcal{M} (2 \cos \theta \vec{u}_r + \sin \theta \vec{u}_\theta)$$

Ce résultat pourrait s'établir par un développement limité comme annoncé précédemment.

3. Complément : action d'un champ magnétique extérieur sur un dipôle magnétique

En guise d'approfondissement, on s'intéresse maintenant aux actions subies par un dipôle magnétique dans un champ magnétique extérieur.

3.1 Résultante exercée sur un dipôle

On part de la force élémentaire de Laplace qui s'exerce sur la boucle de courant :

$$d\vec{F}(M) = Id\vec{OM} \wedge \vec{B}(M)$$

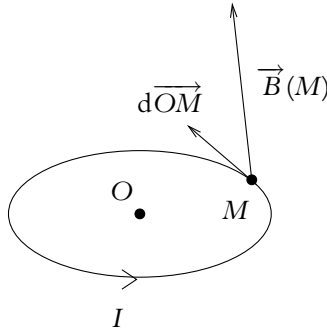


Figure 51.5 Résultante s'exerçant sur un dipôle.

La résultante s'obtient en intégrant sur la boucle :

$$\vec{F} = \oint_{M \in \mathcal{C}} Id\vec{OM} \wedge \vec{B}(M)$$

Si le champ est uniforme, on a :

$$\vec{F} = I \left(\oint_{M \in \mathcal{C}} d\vec{OM} \right) \wedge \vec{B}$$

car le courant I est constant et le champ uniforme. Or $\oint_{M \in \mathcal{C}} d\vec{OM} = \vec{0}$ donc la résultante est nulle :

$$\vec{F} = \vec{0}$$

► **Remarque** Si le champ n'est pas uniforme, on peut établir pour un dipôle rigide :

$$\vec{F} = \text{grad} (\vec{\mathcal{M}} \cdot \vec{B})$$

3.2 Moment exercé sur un dipôle

On peut établir que le moment qui s'exerce sur un dipôle magnétique placé dans un champ magnétique est :

$$\vec{\Gamma}(M) = \vec{\mathcal{M}} \wedge \vec{B}(M)$$

Cette expression est analogue à celle obtenue pour un dipôle électrostatique.

3.3 Énergie potentielle d'un dipôle dans un champ magnétique

Par analogie avec le dipôle électrostatique, on obtient l'expression de l'énergie potentielle d'un dipôle dans un champ magnétique :

$$E_p = -\vec{\mathcal{M}} \cdot \vec{B}$$

De cette expression, on peut déduire l'expression de la force subie par un dipôle rigide (il n'y a pas de déformation). On a donc :

$$\vec{F} = \overrightarrow{\text{grad}} \left(\vec{\mathcal{M}} \cdot \vec{B} \right)$$

Les différentes expressions sont analogues à celles du dipôle électrostatique.

Ces résultats sont valables pour les aimants d'après la modélisation des aimants proposée plus haut.

A. Applications directes du cours

1. Équilibre d'un aimant en présence d'un fil parcouru par un courant

Un aimant peut se déplacer sans frottement dans un plan horizontal. Un fil rectiligne parcouru par un courant d'intensité I est maintenu horizontalement. L'aimant a-t-il des positions d'équilibre ? Si oui, sont-elles stables ?

2. Interaction entre deux aimants

Deux aimants sont maintenus à faible distance l'un de l'autre sur des pivots.

1. Déterminer les positions d'équilibre si les deux aimants sont libres sur leur pivot. Préciser lesquelles sont stables.
2. On bloque un des deux aimants. Indiquer les positions d'équilibre stables de l'aimant mobile (en B) lorsque le pôle Nord de l'aimant fixe (en A) est orienté vers les positions 1, 2, 3, 4 ou 5 indiquées sur la figure suivante :

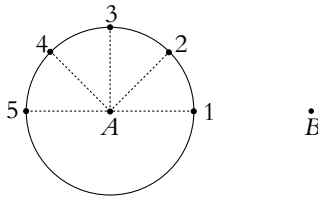


Figure 51.6

B. Exercices et problèmes

1. Interaction entre un aimant et une spire

Un dipôle magnétique de moment $\vec{M}_1 = M_1 \vec{u}_z$ est situé au point O_1 . Une spire circulaire S_2 , de rayon R_2 , parcourue par un courant i_2 constant, a son centre O_2 situé sur l'axe $O_1 z$. La distance $O_1 O_2$ est égale à d , l'axe de la spire est selon $O_1 z$. On considère que $R_2 \ll d$.

1. Déterminer la force de Laplace exercée par le dipôle sur la spire.
2. En utilisant le moment magnétique de la spire S_2 , retrouver l'expression précédente de la force exercée par le dipôle sur S_2 .

2. Champ magnétique terrestre

On suppose dans cet exercice que l'axe des pôles magnétiques est confondu avec l'axe des pôles géographiques. On caractérise alors le champ magnétique terrestre $\vec{B}$ par sa norme et son inclinaison I définie comme l'angle entre $\vec{B}$ et le plan horizontal au point M de la Terre considéré. Gauss a montré que le champ magnétique terrestre peut en première approximation être assimilé au champ magnétique créé par un dipôle magnétique de moment $\vec{M}$ placé au centre O de la Terre.

1. Exprimer $\vec{B}$ en fonction de la latitude λ du point M .
2. En déduire l'inclinaison I en fonction de la latitude.
3. À Paris ($\lambda = 49^\circ$ nord), on mesure une composante horizontale du champ magnétique de $2,05 \cdot 10^{-5}$ T. En déduire la norme du moment $\vec{\mathcal{M}}$ et l'inclinaison I en ce point.
4. Commenter la valeur obtenue sachant qu'on mesure en réalité une inclinaison I de 64° .

Comparaison des propriétés des champs électrostatique et magnétique

52

1. Définitions des champs

Les champs électrostatique $\vec{E}$ et magnétique $\vec{B}$ sont définis à partir de leur action sur une charge ponctuelle q animée d'une vitesse $\vec{v}$ (ce dernier point n'est nécessaire que pour le champ magnétique).

Le champ électrostatique est un vrai vecteur ou vecteur polaire tandis que le champ magnétique est un pseudo-vecteur ou vecteur axial.

2. Flux

Le flux du champ électrostatique à travers une surface fermée est, d'après le théorème de Gauss, égal au rapport de la charge totale située à l'intérieur de cette surface par ϵ_0 .

$$\Phi(\vec{E}) = \frac{Q_{\text{int}}}{\epsilon_0}$$

Le champ magnétique est à flux conservatif à savoir que le flux du champ magnétique à travers une surface fermée est nul.

$$\Phi(\vec{B}) = 0$$

3. Circulation

Le champ électrostatique est à circulation conservative à savoir que la circulation du champ électrostatique le long d'une courbe fermée est nulle : il existe donc un potentiel électrostatique V tel que $\vec{E} = -\text{grad}V$.

$$\oint \vec{E} \cdot d\vec{l} = 0$$

La circulation du champ magnétique le long d'une courbe fermée est, d'après le théorème d'Ampère, égal au produit de μ_0 par l'intensité des courants enlacés par la courbe.

$$\oint \vec{B} \cdot d\vec{l} = \mu_0 I_{\text{enlacé}}$$

On en déduit que :

- les lignes de champ électrique divergent d'une charge positive et convergent vers une charge négative, charges qui créent le champ électrique,
- les lignes de champ magnétique tournent autour des courants qui le créent.

4. Expression à partir de ses sources

Le champ électrostatique s'exprime à partir de ses sources, les charges dq_P au point P , par la loi de Coulomb :

$$d\vec{E}_P(M) = \frac{1}{4\pi\epsilon_0} \frac{dq_P}{(PM)^3} \vec{PM}$$

avec dq_P égal à $\rho(P)d\tau_P$ pour une distribution volumique, $\sigma(P)dS_P$ pour une distribution surfacique ou $\lambda(P)dl_P$ pour une distribution linéique.

Le champ magnétique s'exprime à partir de ses sources, les courants $d\vec{C}_P = I_P d\vec{l}_P$ au point P , par la loi de Biot et Savart :

$$d\vec{B}_P(M) = \frac{\mu_0}{4\pi} \frac{d\vec{C}_P \wedge \vec{PM}}{(PM)^3}$$

5. Propriétés de symétrie

D'après le principe de Curie, les champs électrostatique et magnétique ont les mêmes propriétés de symétrie vis-à-vis des translations et des rotations que les distributions de charges et de courants qui en sont les sources.

5.1 Distributions ayant un plan de symétrie

Le champ électrostatique est symétrique par rapport à un plan de symétrie ; en un point du plan de symétrie, le champ électrostatique appartient à ce plan et les lignes du champ électrostatique sont tangentes au plan de symétrie.

Le champ magnétique est transformé en l'opposé de son symétrique par une symétrie par rapport à un plan de symétrie ; en un point du plan de symétrie, le champ magnétique est perpendiculaire à ce plan et les lignes de champ sont orthogonales au plan de symétrie.

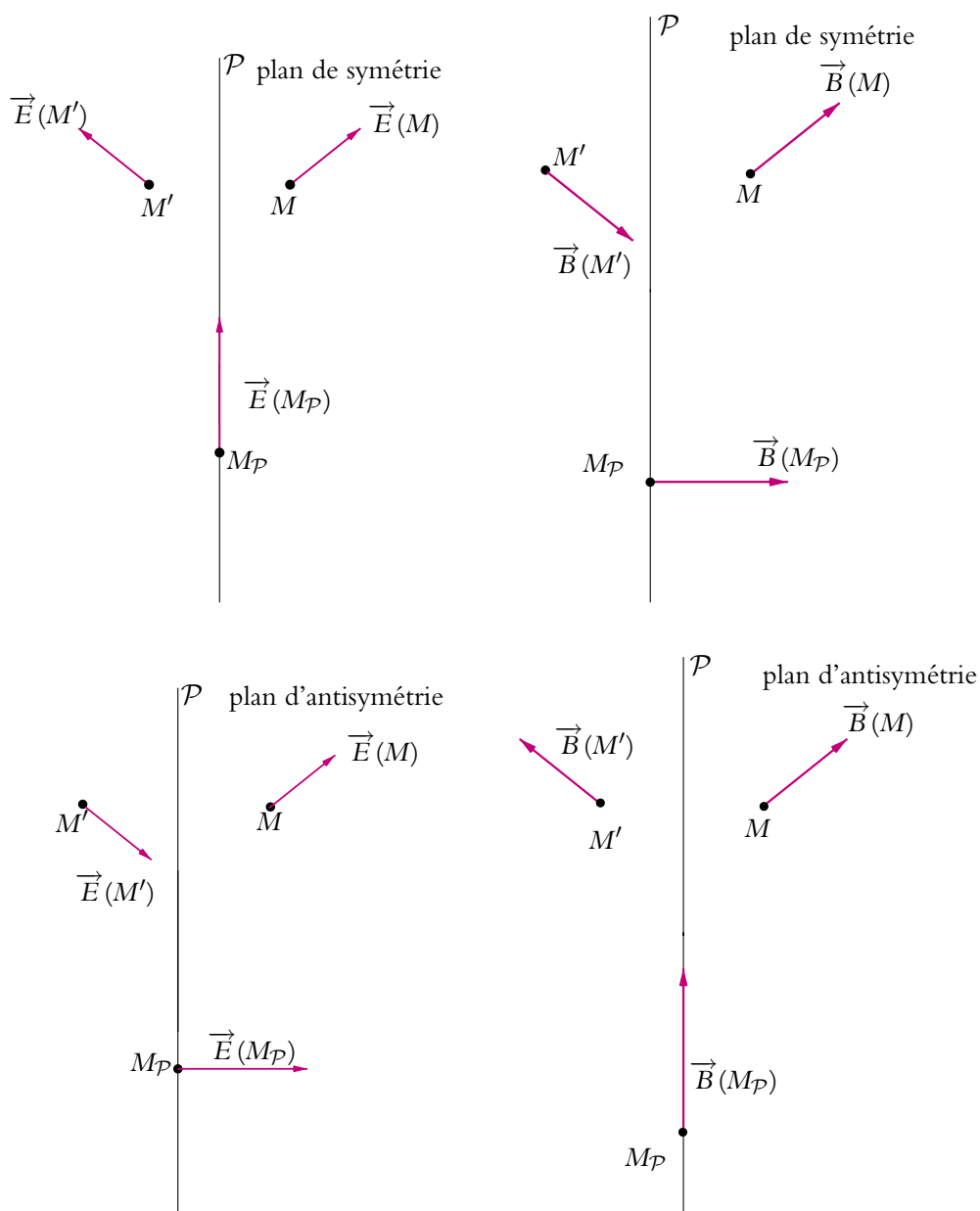


Figure 52.1 Plans de symétrie et d'antisymétrie pour les champs électrostatique et magnétique.

5.2 Distributions ayant un plan d’antisymétrie

Le champ électrostatique est antisymétrique par rapport à un plan d’antisymétrie (le champ électrostatique en un point M' symétrique de M par rapport au plan d’antisymétrie est égal à l’opposé du symétrique par rapport au plan d’antisymétrie du champ électrostatique en M) ; en un point du plan d’antisymétrie, le champ électrostatique est perpendiculaire à ce plan et les lignes du champ électrostatique sont orthogonales au plan d’antisymétrie.

Le champ magnétique est transformé en son symétrique par rapport à un plan de symétrie ; en un point du plan d’antisymétrie, le champ magnétique appartient à ce plan et les lignes de champ sont tangentes au plan d’antisymétrie.

6. Récapitulatif des propriétés des champs électrostatique et magnétique

	Électrostatique	Magnétostatique
Définition force	$\vec{F} = q \vec{E}$	$\vec{F} = q \vec{v} \wedge \vec{B}$
sources	densité de charges ρ (scalaire)	densité de courants $\vec{j}$ (vectorielle)
expression	$\vec{E}(M) = \frac{1}{4\pi\epsilon_0} \iiint \rho(P) d\tau_P \frac{\overrightarrow{PM}}{PM^3}$	$\vec{B}(M) = \frac{\mu_0}{4\pi} \int I_p d\vec{l}_p \wedge \frac{\overrightarrow{PM}}{(PM)^3}$ loi de Biot et Savart
constante	$\frac{1}{\epsilon_0}$	μ_0
caractère	polaire	axial ou pseudo-vecteur
Potentiel	$\vec{E} = -\overrightarrow{\text{grad}} V$	
expression	$V = \frac{1}{4\pi\epsilon_0} \iiint \rho(P) \frac{d\tau_P}{PM}$	
Circulation	$\oint \vec{E}(M) \cdot d\vec{l}_M = 0$ circulation conservative	$\oint \vec{B}(M) \cdot d\vec{l}_M = \mu_0 I_{\text{enlacé}}$ théorème d’Ampère (lien avec les sources)
Flux	$\oiint \vec{E}(M) \cdot d\vec{S}_M = \frac{Q_{\text{int}}}{\epsilon_0}$ théorème de Gauss (lien avec les sources)	$\oiint \vec{B}(M) \cdot d\vec{S}_M = 0$ flux conservatif

7. Dipôles électrostatique et magnétique

	Électrostatique	Magnétostatique
Moment	$\vec{p} = q\vec{NP}$	$\vec{\mathcal{M}} = I\vec{S}$
Potentiel	$V = \frac{1}{4\pi\epsilon_0} \frac{\vec{p} \cdot \vec{OM}}{OM^3}$	
Champ	$\vec{E}(M) = \frac{3(\vec{p} \cdot \vec{OM})\vec{OM} - OM^2\vec{p}}{4\pi\epsilon_0 OM^5}$ $\vec{E} = \begin{cases} E_r = \frac{2p \cos \theta}{4\pi\epsilon_0 r^3} \\ E_\theta = \frac{p \sin \theta}{4\pi\epsilon_0 r^3} \end{cases}$	$\vec{B}(M) = \frac{\mu_0}{4\pi} \frac{3(\vec{\mathcal{M}} \cdot \vec{OM})\vec{OM} - OM^2\vec{\mathcal{M}}}{OM^5}$ $\vec{B} = \begin{cases} B_r = \frac{2\mu_0 \mathcal{M} \cos \theta}{4\pi r^3} \\ B_\theta = \frac{\mu_0 \mathcal{M} \sin \theta}{4\pi r^3} \end{cases}$
Résultante des forces cas du dipôle rigide	$\vec{F} = (\vec{p} \cdot \vec{\text{grad}}) \vec{E}$ $\vec{F} = \vec{\text{grad}} (\vec{p} \cdot \vec{E})$ à p constant	(1) $\vec{F} = \vec{\text{grad}} (\vec{\mathcal{M}} \cdot \vec{B})$ à $\mathcal{M}$ constant
Couple	$\vec{\Gamma} = \vec{p} \wedge \vec{E}$	$\vec{\Gamma} = \vec{\mathcal{M}} \wedge \vec{B}$
Énergie potentielle	$E_p = -\vec{p} \cdot \vec{E}$	$E_p = -\vec{\mathcal{M}} \cdot \vec{B}$

(1) La formule $\vec{F} = (\vec{\mathcal{M}} \cdot \vec{\text{grad}}) \vec{B}$ n'est valable que si le champ magnétique vérifie certaines propriétés, elle n'est pas générale. C'est la seule rubrique où on ne peut pas faire le parallèle avec le dipôle électrostatique.

53

Mouvement d'une particule chargée dans un champ électrique ou magnétique

Le mouvement des particules chargées dans un champ électrique et/ou magnétique est particulièrement important du fait du grand nombre d'applications allant de l'oscilloscope à la télévision en passant par les accélérateurs de particules.

On se limite ici à un traitement classique en excluant tout développement relativiste : la vitesse des particules doit, pour que ce traitement soit justifié, être négligeable devant la vitesse de la lumière à savoir $c = 3.10^8 \text{ m.s}^{-1}$. Or on constatera lors des applications numériques que la vitesse des particules peut assez facilement atteindre des vitesses proches de celle de la lumière. Cette hypothèse pourra alors être évoquée pour expliquer des écarts entre les résultats obtenus théoriquement dans l'approximation classique et les observations expérimentales.

D'autre part, on se limitera aux champs uniformes et indépendants du temps : cela suppose que les champs ne dépendent ni de la position dans l'espace ni de l'instant considérés.

1. Force de Lorentz

1.1 Rappel de l'expression

Comme cela a été vu précédemment, une particule chargée de masse m , de charge q et animée d'une vitesse $\vec{v}$ subit :

- dans un champ électrique une force : $\vec{f} = q\vec{E}$,
- dans un champ magnétique une force : $\vec{f} = q\vec{v} \wedge \vec{B}$.

La force totale en présence d'un champ électrique et d'un champ magnétique est la somme des deux forces précédentes. On appelle *force de Lorentz* la force à laquelle est soumise une particule de masse m , de charge q et animée d'une vitesse $\vec{v}$ dans un

champ électrique et dans un champ magnétique ; elle a pour expression :

$$\vec{f} = q \left(\vec{E} + \vec{v} \wedge \vec{B} \right)$$

On note que $\vec{E}$, $\vec{B}$ et $\vec{v}$ dépendent du référentiel.

1.2 Ordre de grandeur et conséquences

Pour avoir un ordre de grandeur de cette force et de ses composantes électrique et magnétique, on considère le cas d'un électron dont la charge vaut $q = -e = -1,6 \cdot 10^{-19}$ C, animé d'une vitesse égale au dixième de la vitesse de la lumière, à savoir $3 \cdot 10^7$ m.s⁻¹. Cette vitesse correspondra à la limite supérieure admissible pour qu'on puisse négliger les effets relativistes. En effet, dans l'approximation classique, on considère que la vitesse de la particule est négligeable devant la vitesse de la lumière. Admettre que la vitesse puisse valoir un dixième de la vitesse de la lumière revient à accepter 10 % d'erreur, ce qui est le maximum tolérable.

Si on envisage une valeur du champ magnétique de l'ordre de 0,1 T qui est une valeur usuelle de champ magnétique, le module de la composante magnétique de la force de Lorentz sera de l'ordre de :

$$F_{\text{mag}} = qvB = 1,6 \cdot 10^{-19} \times 3 \cdot 10^7 \times 0,1 = 5 \cdot 10^{-13} \text{ N}$$

Pour obtenir une composante électrique de la force de Lorentz du même ordre de grandeur, il faut un champ électrique dont le module soit :

$$E = \frac{F_{\text{élec}}}{q} \simeq \frac{F_{\text{mag}}}{q} = \frac{5 \cdot 10^{-13}}{1,6 \cdot 10^{-19}} = 3 \cdot 10^6 \text{ V.m}^{-1}$$

Il s'agit d'un fort champ électrique. Or c'est la valeur qui donne le même ordre de grandeur que pour un champ magnétique usuel : on en déduit donc que la force d'origine magnétique est généralement plus importante que la force d'origine électrique. On verra notamment la conséquence pratique de ce résultat lors de l'étude du principe de fonctionnement de la télévision.

Il est également important de comparer la force de Lorentz au poids de la particule.

En reprenant le cas de l'électron, sa masse est de l'ordre de 10^{-30} kg. Sur la Terre, qui sera le cas usuel, le champ de pesanteur g est de l'ordre de 10 m.s⁻². L'ordre de grandeur du poids de l'électron est donc de $m_e g = 10^{-30} \times 10 = 10^{-29}$ N, valeur très inférieure à l'ordre de grandeur de la force de Lorentz. Par exemple, dans un champ électrique de l'ordre de 1 kV.m⁻¹, la force électrique vaut $1,6 \cdot 10^{-16}$ N, cette valeur est très grande devant celle du poids.

Pour le proton, l’ordre de grandeur de sa masse est $2 \cdot 10^{-27}$ kg (soit environ 2 000 fois la masse de l’électron) et son poids sur Terre : $m_p g = 2 \cdot 10^{-27} \times 10 = 2 \cdot 10^{-26}$ N, valeur également très inférieure à l’ordre de grandeur de la force de Lorentz (qui est la même pour le proton et l’électron).

Dans la suite, on négligera toujours le poids de la particule chargée.

2. Mouvement dans un champ électrique

On s’intéresse dans un premier temps au mouvement d’une particule chargée dans un champ électrique seul.

2.1 Équation du mouvement

On effectue l’étude en quatre temps de tout problème de mécanique.

- Système : particule chargée de masse m , de vitesse initiale $\vec{v}_0$, de charge q .
- Référentiel : on choisit un référentiel galiléen par exemple le référentiel du laboratoire.
- Bilan des forces :
 - force d’origine électrique : $\vec{f} = q\vec{E}$,
 - poids négligé.
- Choix de la méthode de résolution : on va utiliser le principe fondamental de la dynamique.

Ce dernier s’écrit (en supposant que la masse de la particule ne varie pas au cours du temps) :

$$m\vec{a} = q\vec{E}$$

en notant $\vec{a}$ l’accélération de la particule.

2.2 Étude de la trajectoire

Comme le champ $\vec{E}$ est indépendant du temps, on peut intégrer vectoriellement soit :

$$\vec{v} = \frac{q\vec{E}}{m}t + \vec{v}_0$$

et :

$$\vec{OM} = \frac{1}{2} \frac{q}{m} \vec{E} t^2 + \vec{v}_0 t \quad (53.1)$$

en prenant comme origine O la position initiale de la particule.

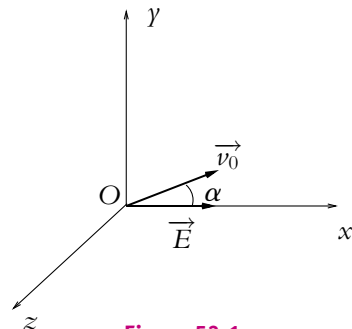


Figure 53.1

On projette l'équation (53.1) dans la base des coordonnées cartésiennes. On choisit l'axe Ox selon le champ électrique et orienté dans le même sens. La vitesse initiale définit avec le champ électrique un plan qu'on prendra comme plan xOy et on notera α l'angle que fait le vecteur vitesse avec le champ électrique dans ce plan. L'axe Oz complétera la base directe.

Par projection dans cette base, on obtient :

$$\begin{cases} x = \frac{qE}{2m}t^2 + v_0 t \cos \alpha \\ y = v_0 t \sin \alpha \\ z = 0 \end{cases}$$

Le mouvement a lieu dans le plan défini par la position initiale, le champ électrique et le vecteur vitesse initial. Ceci n'a rien de surprenant : la seule force qui s'exerce sur la particule est dans ce plan. On remarque l'analogie avec le mouvement d'un point matériel dans le champ de pesanteur : dans les deux cas, la force appliquée est constante et uniforme.

La trajectoire s'obtient en éliminant le temps t dans les relations précédentes ; le plus simple est d'obtenir t en fonction de y . On doit donc distinguer deux cas : le cas où $\alpha = 0$ et le cas où $\alpha \neq 0$, c'est-à-dire le cas où la vitesse initiale est parallèle au champ électrique et celui où elle ne lui est pas parallèle.

a) Trajectoire lorsque la vitesse initiale est parallèle au champ

Dans ce cas, on a $\alpha = 0$ donc $\sin \alpha = 0$ et $\cos \alpha = 1$. La loi horaire du mouvement devient :

$$\begin{cases} x = \frac{qE}{2m}t^2 + v_0 t \\ y = 0 \\ z = 0 \end{cases}$$

On a donc un mouvement rectiligne le long de l'axe Ox .

b) Trajectoire lorsque la vitesse initiale n'est pas parallèle au champ

On peut alors diviser par $v_0 \sin \alpha$ et obtenir : $t = \frac{y}{v_0 \sin \alpha}$. En reportant dans l'équation horaire de x , on en déduit :

$$x = \frac{qE}{2mv_0^2 \sin^2 \alpha} y^2 + y \cotan \alpha$$

Il s'agit d'une parabole passant par l'origine.

On peut chercher l’existence d’un extremum de x en cherchant les valeurs de y pour lesquelles la dérivée $\frac{dx}{dy}$ s’annule :

$$\frac{dx}{dy} = \frac{qE}{mv_0^2 \sin^2 \alpha} y + \cotan \alpha = 0$$

soit pour $y_{\text{ext}} = -\frac{mv_0^2 \sin \alpha \cos \alpha}{qE}$. La valeur correspondante de l’extremum est (en reportant l’expression de y obtenue dans l’équation de la trajectoire) :

$$\begin{aligned} x_{\text{ext}} &= \frac{qE}{2mv_0^2 \sin^2 \alpha} \frac{m^2 v_0^4 \sin^2 \alpha \cos^2 \alpha}{q^2 E^2} - \frac{mv_0^2 \sin \alpha \cos \alpha}{qE} \cotan \alpha \\ &= \frac{mv_0^2 \cos^2 \alpha}{2qE} - \frac{mv_0^2 \cos^2 \alpha}{qE} \quad \text{soit} \quad x_{\text{ext}} = -\frac{mv_0^2 \cos^2 \alpha}{2qE} \end{aligned}$$

Les trajectoires obtenues pour une charge positive ou négative sont représentées sur la figure suivante :

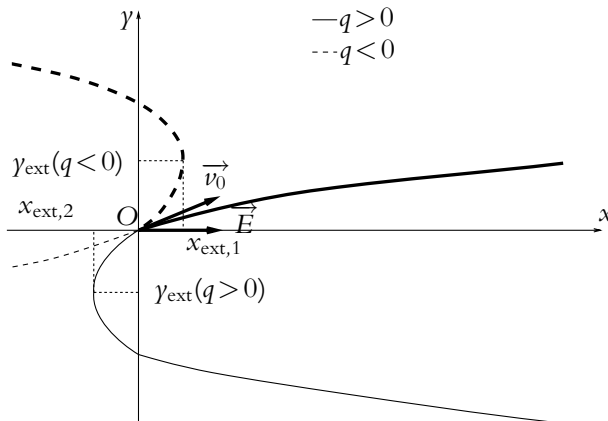


Figure 53.2 Trajectoires d’une particule chargée dans un champ électrique quand la vitesse initiale n’est pas colinéaire au champ.

On retrouve le même type de mouvement que pour un point matériel de masse m dans un champ de pesanteur $\vec{g}$ uniforme.

2.3 Accélération d’une particule chargée par un champ électrique

a) Étude générale

On s’intéresse dans ce paragraphe aux aspects énergétiques afin de déterminer si la vitesse des particules chargées peut être modifiée par action d’un champ électrique.

Le champ électrique $\vec{E}$ est supposé uniforme : le potentiel électrostatique peut s'écrire sous la forme :

$$V = -Ex + \text{constante}$$

De plus, l'énergie potentielle d'une charge q dans ce champ électrostatique s'exprime par :

$$Ep = qV$$

En l'absence de toute autre force que la force électrique, le théorème de l'énergie mécanique donne :

$$\Delta(Ec + Ep) = 0$$

On en déduit la variation d'énergie cinétique :

$$\Delta Ec = -q(V(x_2) - V(x_1))$$

en indiquant par 1 la situation initiale et par 2 la situation finale. Alors, en tenant compte de l'expression du potentiel électrostatique en fonction du champ :

$$\Delta Ec = qE(x_2 - x_1) = \frac{q^2 E^2}{2m} (t_2^2 - t_1^2) + qEv_0 \cos \alpha (t_2 - t_1)$$

- Si $\Delta Ec > 0$, alors l'énergie cinétique de la particule augmente : la particule est accélérée. Ce sera le cas pour les situations suivantes :
 - si $\Delta Ec = \frac{q^2 E^2}{2m} (t_2^2 - t_1^2) > 0$, à savoir si $v_0 = 0$ ou si $\vec{v}_0$ et $\vec{E}$ sont perpendiculaires (ce qui correspond à $\alpha = \pm \frac{\pi}{2}$ soit $\cos \alpha = 0$),
 - si $q \cos \alpha > 0$; cela correspond à deux situations possibles : une charge positive dont la projection de la vitesse initiale selon la direction du champ électrique est dans le même sens que ce dernier ou bien une charge négative dont la projection de la vitesse initiale selon la direction du champ est dans le sens opposé à ce dernier.
- Si $q \cos \alpha < 0$ (c'est-à-dire si une charge positive a une vitesse selon la direction du champ dans le sens opposé à ce dernier ou si une charge négative a une vitesse selon la direction du champ dans le même sens que ce dernier), le second terme est négatif. Le signe de la variation d'énergie cinétique est donc liée aux valeurs absolues de $\frac{qE}{2m}$ et de $v_0 \cos \alpha$.
 - Si $\left| \frac{qE}{2m} (t_2 + t_1) \right| > |v_0 \cos \alpha|$, $\Delta Ec > 0$, l'énergie cinétique augmente : la particule est accélérée.
 - Si $\left| \frac{qE}{2m} (t_2 + t_1) \right| = |v_0 \cos \alpha|$, $\Delta Ec = 0$, l'énergie cinétique est constante et il n'y a ni accélération ni décélération.

- Si $\left| \frac{qE}{2m} (t_2 + t_1) \right| < |\nu_0 \cos \alpha|$, $\Delta Ec < 0$, l’énergie cinétique diminue et la particule est décélérée.

Donc, si on souhaite accélérer une particule, il faut se placer dans l’une des situations suivantes : pas de vitesse initiale, vitesse initiale perpendiculaire au champ, vitesse initiale le long de l’axe du champ dans le même sens que le champ pour une charge positive et dans un sens opposé pour une charge négative ou $\left| \frac{qE}{2m} (t_2 + t_1) \right| > |\nu_0 \cos \alpha|$. On verra dans le paragraphe suivant que si on souhaite que la particule ne soit pas déviée mais seulement accélérée, la seule solution est une vitesse initiale nulle ou une vitesse initiale parallèle au champ électrique.

b) Cas particulier où la vitesse initiale est nulle

Dans ce cas, le mouvement est rectiligne et on a : $x = \frac{qE}{2m} t^2$. Le travail de la force électrique vaut :

$$\int_0^\tau qE \dot{x} dt = \frac{q^2 E^2 \tau^2}{2m} > 0$$

La particule est bien accélérée.

Pour obtenir rapidement la vitesse finale en fonction de U , la différence de potentiel appliquée, on utilise le théorème de l’énergie cinétique :

$$\Delta Ec = \frac{1}{2} m v^2 = -qU$$

On remarque que qU est négatif donc :

$$v = \sqrt{\frac{-2qU}{m}}$$

Exemple d’un proton

$q = 1,6 \cdot 10^{-19}$ C, $m_p = 1,67 \cdot 10^{-27}$ kg et si $U = 1\,000$ V :

$$v = \sqrt{\frac{2 \times 1,6 \cdot 10^{-19} \times 1\,000}{1,67 \cdot 10^{-27}}} = 4,38 \cdot 10^5 \text{ m.s}^{-1}$$

Exemple d’un électron

$q = -1,6 \cdot 10^{-19}$ C, $m_e = 9,1 \cdot 10^{-31}$ kg avec la même tension mais dans l’autre sens :

$$v = \sqrt{\frac{2 \times 1,6 \cdot 10^{-19} \times 1\,000}{9,1 \cdot 10^{-31}}} = 1,88 \cdot 10^7 \text{ m.s}^{-1}$$

On atteint ainsi 6,25 % de la valeur de la vitesse de la lumière avec 1 000 V : il est donc très facile d’atteindre des vitesses relativistes en accélérant des particules chargées

usuelles. Cela limite donc notablement les possibilités de traitement non relativistes même si ce seront les seuls envisagés dans le cadre de ce cours. Dans la suite, on se placera dans des situations telles que l'approximation classique sera valable c'est-à-dire qu'on étudiera ou bien le mouvement d'un proton ou bien celui d'un électron soumis à un potentiel U faible.

2.4 Déviation d'une particule chargée par un champ électrique

On se limite au cas où la vitesse initiale est perpendiculaire au champ électrique. Ce point ne limite en rien la généralité de l'étude : la vitesse initiale peut être décomposée en une composante parallèle au champ (qui donne lieu à une accélération) et une composante perpendiculaire qu'on va analyser ici.

D'autre part, on suppose que le champ électrique ne s'applique que dans une zone limitée de l'espace de taille L dans la direction perpendiculaire au champ ($0 < y < L$ en gardant les notations précédentes).

Dans ce cas, les équations horaires du mouvement sont les suivantes :

$$\begin{cases} x = \frac{qE}{2m}t^2 \\ y = v_0 t \end{cases}$$

La particule quitte la zone où règne le champ pour $y_1 = L$ soit à l'instant $t = \frac{L}{v_0}$. Son abscisse est alors :

$$x_1 = \frac{qEL^2}{2mv_0^2}$$

L'angle de déviation de la trajectoire est donné par :

$$\tan \theta = \tan(\widehat{\vec{v}_0, \vec{v}_1}) = \frac{\dot{x}_1}{\dot{y}_1}$$

$$\text{Or } \dot{x} = \frac{qEt}{m} \text{ et } \dot{y} = v_0. \text{ Donc } \dot{x}_1 = \frac{qEL}{mv_0} \text{ et } \tan \theta = \frac{qEL}{mv_0^2}$$

Exemple

$E = 3\,000 \text{ V.m}^{-1}$, $L = 2 \text{ cm}$, $q = 1,6 \cdot 10^{-19} \text{ C}$, $m = 1,67 \cdot 10^{-27} \text{ kg}$
et $v_0 = 4,38 \cdot 10^5 \text{ m.s}^{-1} \ll c$ alors

$$\tan \theta = \frac{1,6 \cdot 10^{-19} \times 3\,000 \times 2 \cdot 10^{-2}}{1,67 \cdot 10^{-27} \times (4,38 \cdot 10^5)^2} = 3,00 \cdot 10^{-2} \quad \text{et} \quad \theta = 1,72^\circ$$

Cet angle est faible : la vitesse est presque parallèle à Oy et le proton est peu dévié.

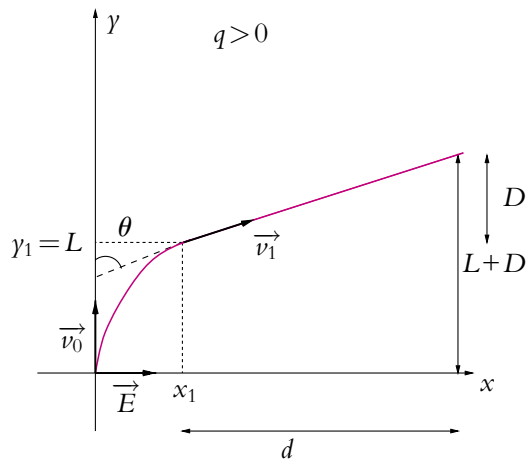


Figure 53.3 Déviation d'une particule chargée par un champ électrique.

On peut également calculer la déviation $\delta = x_1 + d$ sur un écran situé à une distance D des plaques de déviation car la tangente de l'angle de déviation peut aussi s'exprimer en fonction de D et de d :

$$\tan \theta = \frac{d}{D} \quad \text{d'où} \quad \delta = \frac{qEL}{mv_0^2} \left(\frac{L}{2} + D \right)$$

Exemple

$D = 10$ cm pour le cas précédent.

$$\delta = \frac{1,6 \cdot 10^{-19} \times 3000 \times 2 \cdot 10^{-2}}{1,67 \cdot 10^{-27} \times (4,38 \cdot 10^5)^2} \times \left(\frac{2 \cdot 10^{-2}}{2} + 10 \cdot 10^{-2} \right) = 3,3 \text{ mm}$$

2.5 Application à l'oscilloscope

Le principe de fonctionnement des oscilloscopes analogiques repose sur ces deux aspects d'accélération et de déviation par un champ électrique.

La figure ci-dessous représente la constitution interne d'un oscilloscope.

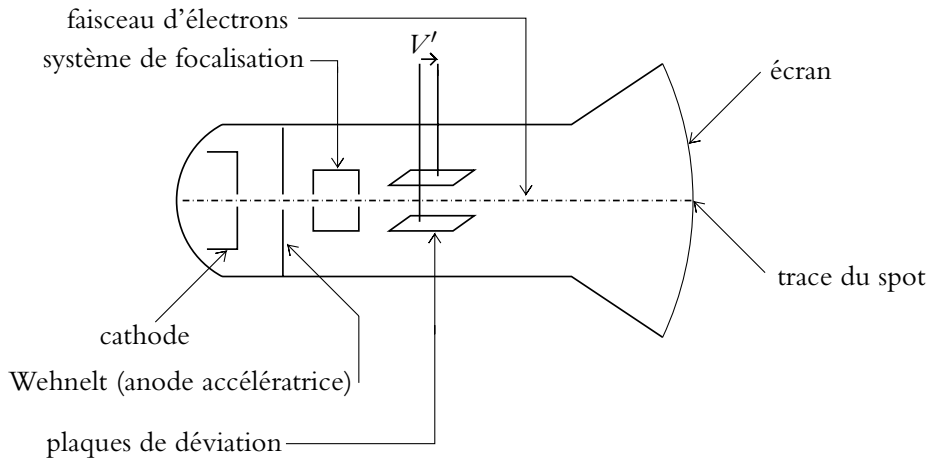


Figure 53.4 Constitution de l'oscilloscope.

Le premier élément est un canon à électrons qui permet de constituer un faisceau d'électrons homocinétiques (c'est-à-dire d'électrons ayant la même vitesse) à l'aide d'une tension accélératrice¹ qui vaut généralement 2 000 V. Les électrons initialement

¹ On l'appelle tension de Wehnelt.

au repos se retrouvent animés d'une vitesse

$$\nu_0 = \sqrt{\frac{2eV}{m}} = \sqrt{\frac{2 \times 1,6 \cdot 10^{-19} \times 2\,000}{9,1 \cdot 10^{-31}}}$$

$$\nu_0 = 26,5 \cdot 10^6 \text{ m.s}^{-1}$$

À la sortie de ce canon à électrons, on a donc un faisceau d'électrons de même vitesse ν_0 .

Les valeurs obtenues pour la vitesse dans le cadre de la mécanique classique sont assez proches de la vitesse de la lumière. Il serait préférable d'adopter un traitement relativiste : on trouverait alors des valeurs plus faibles ($\nu_0 \simeq 11,9 \cdot 10^6 \text{ m.s}^{-1}$). La mécanique classique donne le bon ordre de grandeur et permet d'expliquer le principe de fonctionnement. En revanche, les résultats numériques ne sont pas parfaitement exacts.

On applique alors une tension V' entre deux plaques déflectrices distantes de d telles que le champ électrique qui est perpendiculaire aux plaques soit perpendiculaire à la vitesse des électrons. Ces derniers sont déviés d'un angle θ tel que :

$$\tan \theta = \frac{qEL}{m\nu_0^2}$$

en notant $E = \frac{V'}{d}$ le champ électrique entre les plaques et L la largeur des plaques définissant la zone où s'applique le champ.

Sur l'écran situé à une distance D , le faisceau sera décalé de :

$$x = \frac{qEL}{m\nu_0^2} \left(\frac{L}{2} + D \right)$$

La déviation est donc proportionnelle au champ électrique ou à la différence de potentiel appliquée. On obtient ainsi facilement la valeur de la tension à partir de la connaissance des calibres.

3. Mouvement d'une particule chargée dans un métal - modèle de la loi d'Ohm locale

3.1 Mouvement d'une particule chargée dans un métal

Dans ce paragraphe, on s'intéresse au mouvement des particules chargées au sein d'un métal placé dans un champ électrique. La différence par rapport à ce qui précède vient de la nécessité de prendre en compte l'environnement. Il s'agit par exemple d'étudier le mouvement d'un électron dans un câble électrique de cuivre. Le fait

d'appliquer aux extrémités du câble une différence de potentiel revient à appliquer un champ électrique. La présence des atomes de cuivre et des autres électrons sera prise en compte en introduisant une force de frottement proportionnelle à la vitesse qui traduit l'interaction entre les électrons de conduction et leur environnement (réseau cristallin, autres électrons, etc). Ce modèle de la conduction électrique dans les métaux est appelé *modèle de Drude*.

On reprend l'étude en quatre temps de tout problème de mécanique.

- Système : particule chargée de masse m , de vitesse initiale $\vec{v}_0$, de charge q .
- Référentiel : on choisit un référentiel galiléen par exemple le référentiel du laboratoire.
- Bilan des forces :
 - force d'origine électrique : $\vec{f} = q\vec{E}$,
 - poids négligé,
 - force de frottement : $\vec{f} = -\frac{m}{\tau}\vec{v}$ où τ correspond à l'« âge moyen des porteurs » c'est-à-dire à la durée moyenne entre deux chocs. L'ordre de grandeur de τ est 10^{-14} s.
- Choix de la méthode de résolution : on va utiliser le principe fondamental de la dynamique.

Ce dernier s'écrit :

$$m\vec{a} = q\vec{E} - \frac{m}{\tau}\vec{v}$$

en notant $\vec{a}$ l'accélération de la particule. On obtient l'équation différentielle vectorielle du premier ordre suivante :

$$\frac{d\vec{v}}{dt} + \frac{1}{\tau}\vec{v} = \frac{q}{m}\vec{E}$$

L'équation homogène associée : $\frac{d\vec{v}}{dt} + \frac{1}{\tau}\vec{v} = \vec{0}$ admet pour solution :

$$\vec{v} = \vec{V} \exp\left(-\frac{t}{\tau}\right).$$

On cherche une solution particulière constante de l'équation totale (son second membre est constant). On obtient : $\vec{v} = \frac{q\tau}{m}\vec{E}$.

La solution globale s'écrit :

$$\vec{v} = \vec{V} \exp\left(-\frac{t}{\tau}\right) + \frac{q\tau}{m}\vec{E}$$

Après un régime transitoire dont la durée caractéristique est τ , on atteint une vitesse limite :

$$\vec{v}_{\text{lim}} = \frac{q\tau}{m}\vec{E}$$

3.2 Interprétation de la loi d'Ohm locale

Cette étude permet de donner une interprétation microscopique à la loi d'Ohm locale reliant la densité de courant $\vec{j}$ et le champ électrique $\vec{E}$ appliqué. La densité de courant s'exprime par : $\vec{j} = nq\vec{v}$ où n est le nombre moyen de particules de charge q ayant une vitesse² $\vec{v}$.

D'après l'étude du modèle de Drude, lorsque ces porteurs sont placés dans un champ électrique $\vec{E}$ (ou soumis à une différence de potentiel V telle que $\vec{E} = -\text{grad}V$), ils acquièrent, après un régime transitoire de durée caractéristique τ , une vitesse :

$$\vec{v} = \frac{q\tau}{m} \vec{E}$$

On en déduit l'expression de la densité de courant :

$$\vec{j} = \frac{nq^2\tau}{m} \vec{E}$$

Or la loi d'Ohm locale est donnée par la relation $\vec{j} = \gamma \vec{E}$, relation qui peut être interprétée grâce au modèle précédent. La conductivité γ s'identifie à :

$$\gamma = \frac{nq^2\tau}{m}$$

On notera que γ est indépendant du signe des charges : les effets de charges opposés s'ajoutent. C'est la base de la conductimétrie en chimie.

Exemple du cuivre

On peut calculer le temps caractéristique du régime transitoire à partir des données expérimentales de la conductivité pour le cuivre $\gamma = 5,9 \cdot 10^7 \text{ S.m}^{-1}$.

Le nombre d'électrons conducteurs n pour le cuivre s'obtient en considérant qu'on a un électron de conduction par atome de cuivre. Il suffit donc d'avoir le nombre d'atomes de cuivre contenus dans un mètre cube de cuivre. Ce nombre est égal au rapport de la masse volumique du cuivre ($8,9 \cdot 10^6 \text{ g.m}^{-3}$) par la masse molaire ($63,5 \text{ g.mol}^{-1}$) multiplié par le nombre d'Avogadro ($6,02 \cdot 10^{23}$) soit :

$$n = \frac{8,90 \cdot 10^6}{63,5} \times 6,02 \cdot 10^{23} = 8,43 \cdot 10^{28} \text{ m}^{-3}$$

On rappelle que la masse de l'électron est $m_e = 9,1 \cdot 10^{-31} \text{ kg}$ donc

$$\tau = \frac{m\gamma}{nq^2} = \frac{9,1 \cdot 10^{-31} \times 5,9 \cdot 10^7}{8,43 \cdot 10^{28} \times (1,6 \cdot 10^{-19})^2} = 2,5 \cdot 10^{-14} \text{ s}$$

² On fait l'hypothèse qu'on n'a qu'un seul type de porteurs. Sinon, il faudrait sommer sur l'ensemble des types de porteurs : $\vec{j} = \sum_i n_i q_i \vec{v}_i$ où n_i est le nombre moyen de particules de charge q_i ayant une vitesse $\vec{v}_i$.

On en déduit que le régime permanent est établi presque instantanément. Il faut bien évidemment pour que ceci soit valable que le champ électrique, s'il dépend du temps, ne varie pas plus rapidement, à savoir que la tension appliquée ait une fréquence très inférieure à 10^{14} Hz. Dans le cas contraire, les électrons réagissent avec un certain retard aux variations du champ électrique et le modèle qui vient d'être décrit n'est plus valable.

3.3 Lien avec la loi d'Ohm intégrale

On va établir la loi d'Ohm intégrale $U = RI$ à partir de la loi d'Ohm locale $\vec{j} = \gamma \vec{E}$ dans le cas d'un circuit filiforme.

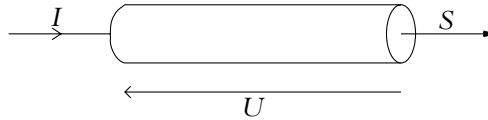


Figure 53.5 Conducteur ohmique.

L'intensité I est obtenue en calculant le flux de $\vec{j}$ à travers la section du fil :

$$I = \iint_{\text{section}} \vec{j} \cdot \vec{dS} = \iint_{\text{section}} \gamma \vec{E} \cdot \vec{dS}$$

soit

$$I = \gamma \iint_{\text{section}} \vec{E} \cdot \vec{dS}$$

La tension ou différence de potentiel se détermine en calculant la circulation du champ électrique le long du fil :

$$U = \int_{\text{fil}} dV = \int_{\text{fil}} \overrightarrow{\text{grad}} V \cdot d\overrightarrow{OM} = - \int_{\text{fil}} \vec{E} \cdot d\overrightarrow{OM}$$

Donc :

$$\frac{U}{I} = \frac{- \int_{\text{fil}} \vec{E} \cdot d\overrightarrow{OM}}{\gamma \iint_{\text{section}} \vec{E} \cdot \vec{dS}}$$

Si on augmente l'intensité du champ électrique d'un facteur λ , on aura :

$$\frac{U'}{I'} = \frac{-\lambda \int_{\text{fil}} \vec{E} \cdot d\overrightarrow{OM}}{\lambda \gamma \iint_{\text{section}} \vec{E} \cdot \vec{dS}} = \frac{- \int_{\text{fil}} \vec{E} \cdot d\overrightarrow{OM}}{\gamma \iint_{\text{section}} \vec{E} \cdot \vec{dS}} = \frac{U}{I}$$

Le rapport $\frac{U}{I}$ est donc constant : on le note R et on le définit comme la résistance du fil. La relation précédente permet de calculer la résistance d'un conducteur de forme quelconque. Le rapport des intégrales donnant la circulation et le flux du champ électrique ne dépend que de la géométrie du circuit et non de l'intensité du champ $\vec{E}$.

Exemple : fil de section constante S et de longueur l

Le champ $\vec{E}$ est appliqué le long du fil : $\vec{E}$ est parallèle à l'axe du fil et $\vec{E}$, $d\vec{S}$ et $d\vec{OM}$ sont colinéaires. On en déduit :

$$I = \gamma \iint_{\text{section}} \vec{E} \cdot d\vec{S} = \gamma ES \quad ; \quad U = - \int_{\text{fil}} \vec{E} \cdot d\vec{OM} = El$$

donc $U = RI$ avec

$$R = \frac{l}{\gamma S}$$

Par exemple, pour un fil de cuivre de longueur $l = 30$ cm, de section $S = 1$ mm², de conductivité $\gamma = 5,9 \cdot 10^7$ S.m⁻¹, on trouve :

$$R = \frac{30 \cdot 10^{-2}}{(10^{-3})^2 \times 5,9 \cdot 10^7} = 5 \cdot 10^{-3} \Omega$$

Cette valeur est très faible devant les valeurs de résistance usuelles, c'est pourquoi on néglige les résistances des fils dans l'étude des circuits électriques.

4. Mouvement dans un champ magnétique

Soit une particule de masse m , de charge q et de vitesse initiale $\vec{v}_0$ placée dans un champ magnétique $\vec{B}$. On étudie le mouvement de cette particule.

4.1 Équation du mouvement - fréquence cyclotron

- Système : la particule.
- Référentiel : terrestre ou du laboratoire considéré comme galiléen.
- Bilan des forces :
 - force magnétique $\vec{f} = q \vec{v} \wedge \vec{B}$,
 - poids qui sera négligé compte tenu des ordres de grandeur.
- On choisit d'appliquer le principe fondamental de la dynamique :

$$m \vec{a} = q \vec{v} \wedge \vec{B}$$

soit

$$\frac{d\vec{v}}{dt} = -\frac{q}{m} \vec{B} \wedge \vec{v} = \vec{\Omega} \wedge \vec{v}$$

en notant que $\vec{\Omega} = -\frac{q}{m} \vec{B}$ est un vecteur constant si $\vec{B}$ l'est.

On note $\omega_c = \left\| \vec{\Omega} \right\| = \frac{|q| B}{m}$. Quelle est la dimension de cette quantité ? La dimension du champ magnétique est :

$$[B] = \frac{[F]}{[Q][v]} = \frac{[M][L][T]^{-2}}{[Q][L][T]^{-1}} = [M][T]^{-1}[Q]^{-1}$$

On en déduit que :

$$[\Omega] = \frac{[Q][B]}{[M]} = \frac{1}{[T]}$$

$\left\| \vec{\Omega} \right\|$ est donc homogène à l'inverse d'un temps, c'est-à-dire à une fréquence ou une pulsation.

Le module ω_c de $\vec{\Omega}$ est appelé *pulsation cyclotron*.

Dans ce cas, l'intégration de l'équation du mouvement donne :

$$\vec{v} = \vec{\Omega} \wedge \vec{OM} + \vec{v}_0 \quad (53.2)$$

car $\vec{\Omega}$ est constant.

4.2 Travail de la force magnétique

Le travail élémentaire de la force magnétique lors d'un déplacement élémentaire de la particule située en M s'écrit :

$$\delta W = (q \vec{v} \wedge \vec{B}) \cdot d\vec{OM}$$

Or $d\vec{OM} = \vec{v} dt$ est colinéaire à $\vec{v}$ et donc perpendiculaire à $q \vec{v} \wedge \vec{B}$. On en déduit :

$$\delta W = 0 \quad \text{et} \quad W = 0$$

La force magnétique ne travaille pas.

Les conséquences de ce résultat sont importantes. En effet, l'application du théorème de l'énergie cinétique à la particule permet de conclure à la conservation de l'énergie cinétique car le travail de la force est nul.

L'énergie cinétique d'une particule chargée dans un champ magnétique seul est une constante du mouvement.

La masse de la particule ne changeant pas et l'énergie cinétique s'exprimant par $E_c = \frac{1}{2}mv^2$, on en déduit que **le module du vecteur vitesse de la particule est constant au cours du mouvement**.

4.3 Étude de la trajectoire

a) Angle entre le vecteur vitesse et le champ magnétique

Le vecteur vitesse peut se décomposer comme la somme de deux composantes l'une $\vec{v}_{\parallel}$ étant parallèle au champ magnétique $\vec{B}$ et l'autre $\vec{v}_{\perp}$ lui étant orthogonale :

$$\vec{v} = \vec{v}_{\parallel} + \vec{v}_{\perp} \quad \text{avec} \quad \vec{v}_{\parallel} = \left(\vec{v} \cdot \frac{\vec{B}}{\|\vec{B}\|} \right) \frac{\vec{B}}{\|\vec{B}\|^2}$$

Or :

$$\frac{d}{dt} \left(\vec{v} \cdot \vec{B} \right) = \frac{d\vec{v}}{dt} \cdot \vec{B} = \frac{q\vec{v} \wedge \vec{B}}{m} \cdot \vec{B} = 0$$

donc $\vec{v} \cdot \vec{B}$ est constant³.

On en déduit que :

1. la composante parallèle $\vec{v}_{\parallel}$ de la vitesse est une constante du mouvement,
2. comme $\|\vec{v}\|$ est constant, le module $v_{\perp}$ de la composante perpendiculaire $\vec{v}_{\perp}$ l'est aussi.
3. l'angle θ entre $\vec{B}$ et $\vec{v}$ est constant, car les modules des deux vecteurs le sont et que :

$$\vec{v} \cdot \vec{B} = vB \cos \theta = \text{constante}$$

b) Projection de la trajectoire dans le plan perpendiculaire au champ magnétique

La projection de l'équation du mouvement sur le plan perpendiculaire au champ magnétique s'écrit (l'indice $\perp$ désigne les grandeurs dans ce plan) :

$$m \frac{d\vec{v}_{\perp}}{dt} = q\vec{v}_{\perp} \wedge \vec{B}$$

soit

$$\frac{d\vec{v}_{\perp}}{dt} = \vec{\Omega} \wedge \vec{v}_{\perp}$$

avec $\vec{\Omega} = -\frac{q}{m} \vec{B}$ qui est un vecteur constant.

Or, pour un mouvement circulaire, on peut écrire la vitesse sous la forme : $\vec{v} = \vec{\Omega}' \wedge \vec{OM}$ en notant $\vec{\Omega}'$ le vecteur rotation. On en déduit l'expression de

³ On peut également établir ce résultat en utilisant la relation (53.2) :

$$\vec{v} \cdot \vec{B} = \left(\vec{\Omega}' \wedge \vec{OM} \right) \cdot \vec{B} + \vec{v}_0 \cdot \vec{B} = \left(-\frac{q}{m} \vec{B} \wedge \vec{OM} \right) \cdot \vec{B} + \vec{v}_0 \cdot \vec{B} = \vec{v}_0 \cdot \vec{B}$$

l'accélération pour un mouvement circulaire uniforme :

$$\vec{a} = \vec{\Omega}' \wedge \vec{v}$$

Réciproquement, on peut montrer que si $\vec{a} = \vec{\Omega} \wedge \vec{v}$, le mouvement est circulaire uniforme. On en déduit que la projection de la trajectoire dans le plan perpendiculaire au champ magnétique est circulaire uniforme de vecteur rotation $\vec{\Omega}$.

Soit en projetant suivant la direction radiale des coordonnées cylindriques d'axe parallèle au champ magnétique :

$$mR\Omega^2 = m\frac{v_{\perp}^2}{R} = |q| v_{\perp} B$$

soit un rayon :

$$R = \frac{mv_{\perp}}{|q| B}$$

c) Trajectoire

La trajectoire est donc composée :

- d'un mouvement de translation rectiligne uniforme parallèlement au champ magnétique $\vec{B}$ (car la projection de la vitesse dans la direction du champ magnétique $v_{\parallel}$ est constante),
- d'un mouvement circulaire dans le plan perpendiculaire au champ magnétique $\vec{B}$ de rayon : $R_{\perp} = \frac{mv_{\perp}}{|q| B}$, décrit à la vitesse angulaire $\omega_c = \frac{|q| B}{m}$.

Dans le cas général, il s'agit d'un mouvement hélicoïdal. Le mouvement sera circulaire si la vitesse initiale est perpendiculaire au champ magnétique (soit $v_{\parallel} = 0$).

d) Équations horaires du mouvement

Les équations horaires ne sont pas nécessaires pour déterminer la nature du mouvement puisque les études précédentes l'ont déjà établi. Cependant, il est intéressant de les obtenir : c'est en effet l'occasion de présenter deux méthodes pour intégrer un système d'équations différentielles couplées, utiles dans d'autres cas.

Pour ce faire, on choisit la base de projection suivante : l'axe Oz est pris selon le champ magnétique $\vec{B}$, l'axe Oy perpendiculaire à Oz dans le plan contenant $\vec{B}$ et la vitesse initiale $\vec{v}_0$ qu'on suppose faire un angle θ avec l'axe Oz , l'axe Ox complète la base pour que $Oxyz$ soit directe. L'origine O sera la position initiale de la particule chargée.

Le principe fondamental de la dynamique :

$$m \vec{a} = q \vec{v} \wedge \vec{B}$$

se traduit dans cette base par :

$$m \begin{vmatrix} \ddot{x} \\ \ddot{y} \\ \ddot{z} \end{vmatrix} = q \begin{vmatrix} \dot{x} \\ \dot{y} \\ \dot{z} \end{vmatrix} \wedge \begin{vmatrix} 0 \\ 0 \\ B \end{vmatrix} \quad \text{soit} \quad \begin{cases} m\ddot{x} = qB\dot{y} \\ m\ddot{y} = -qB\dot{x} \\ m\ddot{z} = 0 \end{cases}$$

On commence par intégrer l'équation différentielle en z : $\ddot{z} = 0$ qui donne : $\dot{z} = v_0 \cos \theta = \text{constante}$ puis $z = v_0 \cos \theta t$ en utilisant les conditions initiales.

Il reste maintenant à résoudre le système d'équations différentielles suivant :

$$\begin{cases} m\ddot{x} = qB\dot{y} \\ m\ddot{y} = -qB\dot{x} \end{cases}$$

Ces équations sont couplées. Deux approches sont envisageables :

- intégrer l'une, injecter le résultat dans l'autre avant de la résoudre puis reprendre la première avec la solution connue,
- utiliser une combinaison linéaire complexe et résoudre l'équation obtenue dans l'espace des complexes.

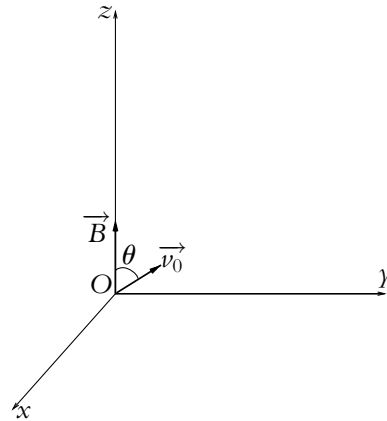


Figure 53.6 Mouvement d'une particule chargée dans un champ magnétique.

On va utiliser successivement ces deux méthodes :

1. Méthode de « substitution » :

On commence⁴ par exemple par intégrer en $\dot{x}$ par rapport au temps :

$$m\dot{x} = qBy + C$$

où C est une constante d'intégration. Or, en tenant compte des conditions initiales : $\dot{x}(0) = 0$ et $y(0) = 0$ donc $C = 0$. On a donc $\dot{x} = \frac{qB}{m}y$ qu'on reporte dans l'autre équation :

$$\ddot{y} + \left(\frac{qB}{m}\right)^2 y = 0 \quad \Leftrightarrow \quad \ddot{y} + \omega_c^2 y = 0$$

dont la solution s'écrit : $y = C_1 \cos(\omega_c t) + C_2 \sin(\omega_c t)$. Les constantes s'obtiennent à partir des conditions initiales : $y(0) = 0 = C_1$ et $\dot{y} = v_0 \sin \theta = \omega_c C_2$ soit :

⁴ On peut éviter cette première étape en utilisant la relation : $\vec{v} = \vec{\Omega} \wedge \vec{OM} + \vec{v}_0$ et en la projetant dans la base utilisée ici.

$$C_1 = 0 \text{ et } C_2 = \frac{\nu_0 \sin \theta}{\omega_c} \text{ et}$$

$$\gamma = \frac{\nu_0 \sin \theta}{\omega_c} \sin(\omega_c t)$$

On reporte alors dans l'équation en $\dot{x}$:

$$\dot{x} = \omega_c \frac{\nu_0 \sin \theta}{\omega_c} \sin(\omega_c t) = \nu_0 \sin \theta \sin(\omega_c t)$$

L'intégration donne : $x = -\frac{\nu_0 \sin \theta}{\omega_c} \cos(\omega_c t) + C$. Or $x(0) = C - \frac{\nu_0 \sin \theta}{\omega_c} = 0$

donc $C = \frac{\nu_0 \sin \theta}{\omega_c}$ et

$$x = \frac{\nu_0 \sin \theta}{\omega_c} (1 - \cos(\omega_c t))$$

2. Méthode par les complexes :

On pose $\underline{X} = x + i\gamma$ en notant i la quantité telle que $i^2 = -1$ soit $x = \text{Re}(\underline{X})$ et $\gamma = \text{Im}(\underline{X})$. Alors, en sommant l'équation en $\ddot{x}$ et i fois l'équation en $\ddot{\gamma}$:

$$\ddot{\underline{X}} = \ddot{x} + i\ddot{\gamma} = \omega_c (\dot{\gamma} - i\dot{x}) = -i\omega_c \dot{\underline{X}}$$

On a donc une équation différentielle complexe du premier ordre en $\underline{Y} = \dot{\underline{X}}$:

$$\dot{\underline{Y}} + i\omega_c \underline{Y} = 0$$

dont la solution s'écrit : $\underline{Y} = \dot{\underline{X}} = A_1 \exp(-i\omega_c t)$.

Les conditions initiales imposent : $\dot{\underline{X}}(0) = \dot{x}(0) + i\dot{\gamma}(0) = i\nu_0 \sin \theta$ et $A_1 = i\nu_0 \sin \theta$.

Il ne reste plus qu'à intégrer :

$$\dot{\underline{X}} = i\nu_0 \sin \theta \exp(-i\omega_c t)$$

en :

$$\underline{X} = -\frac{\nu_0 \sin \theta}{\omega_c} \exp(-i\omega_c t) + A_2$$

À $t = 0$, on a : $\underline{X}(0) = x(0) + i\gamma(0) = 0 = -\frac{\nu_0 \sin \theta}{\omega_c} + A_2$ donc $A_2 = \frac{\nu_0 \sin \theta}{\omega_c}$ et :

$$\underline{X} = \frac{\nu_0 \sin \theta}{\omega_c} (1 - \exp(-i\omega_c t)) = x + i\gamma$$

Par identification des parties réelle et imaginaire, on retrouve les équations horaires :

$$\begin{cases} x = \frac{\nu_0 \sin \theta}{\omega_c} (1 - \cos(\omega_c t)) \\ \gamma = \frac{\nu_0 \sin \theta}{\omega_c} \sin(\omega_c t) \end{cases}$$

On retrouve un mouvement hélicoïdal dans le cas général et un mouvement circulaire quand la vitesse initiale est perpendiculaire au champ magnétique. En effet, on distingue :

- une translation rectiligne uniforme le long de l'axe du champ magnétique $\vec{B}$:

$$z = v_0 \cos \theta t$$

- un mouvement circulaire dans le plan xOy perpendiculaire au champ. On vérifie que :

$$\left(\frac{\omega_c x}{v_0 \sin \theta} - 1 \right)^2 + \left(\frac{\omega_c y}{v_0 \sin \theta} \right)^2 = 1$$

ou :

$$(x - x_0)^2 + y^2 = \rho^2$$

en notant $x_0 = \rho = \frac{v_0 \sin \theta}{\omega_c}$. On obtient donc l'équation d'un cercle de centre $C(x_0, 0)$ et de rayon ρ . Ce cercle est décrit à la vitesse angulaire ω_c . On retrouve bien les mêmes résultats.

e) Mouvement hélicoïdal d'une particule chargée dans un champ magnétique

Il s'agit du cas général où la vitesse initiale n'est pas perpendiculaire au champ magnétique : on observe alors une dérive parallèlement à l'axe du champ et un mouvement circulaire perpendiculairement à l'axe.

La particule décrit donc une hélice dont les paramètres sont les suivantes :

- la fréquence de rotation perpendiculairement à l'axe du champ est donnée par : $f = \frac{\omega_c}{2\pi} = \frac{qB}{2\pi m}$ soit une période :

$$T = \frac{2\pi}{\omega_c} = \frac{2\pi m}{qB}$$

- le pas de l'hélice défini comme la variation de hauteur de la particule quand elle a tourné de 2π ou comme la distance entre deux « spires » le long de l'axe du champ vaut :

$$h = v_0 \cos \theta T = \frac{2\pi v_0 \cos \theta}{\omega_c}$$

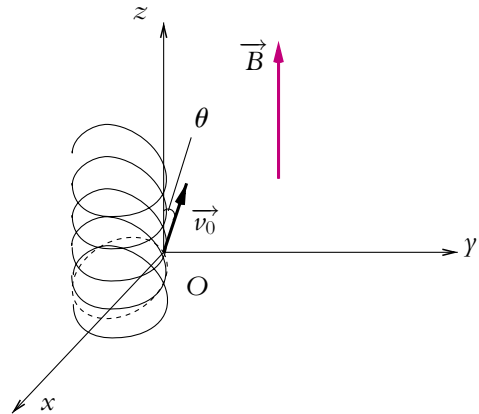


Figure 53.7 Mouvement hélicoïdal d'une particule chargée dans un champ magnétique.

f) Mouvement circulaire d'une particule chargée dans un champ magnétique

Le mouvement se réduit à un mouvement circulaire si la vitesse initiale est perpendiculaire au champ magnétique. Le rayon de la trajectoire est alors :

$$R = \frac{mv_0}{|q|B}$$

On notera que le signe de la charge q change le sens de la trajectoire comme cela est indiqué sur le schéma ci-après :

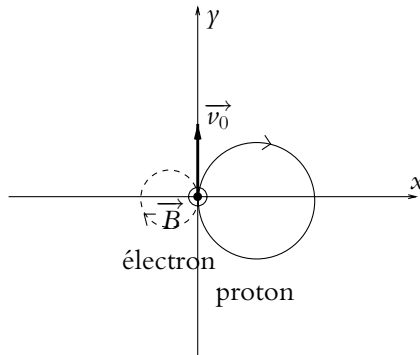


Figure 53.8 Sens de rotation d'une particule chargée dans un champ magnétique en fonction du type de charge.

On considère des particules préalablement accélérées par une différence de potentiel de 1 kV comme cela a été détaillé plus haut. La valeur du rayon de la trajectoire pour un proton de masse $m_p = 1,67 \cdot 10^{-27}$ kg, de charge $q = 1,6 \cdot 10^{-19}$ C, de vitesse $v_0 = 4,38 \cdot 10^5$ m.s⁻¹ et soumis à un champ magnétique de 0,1 T est :

$$R = \frac{1,67 \cdot 10^{-27} \times 4,38 \cdot 10^5}{1,6 \cdot 10^{-19} \times 0,1} = 4,57 \text{ cm}$$

et pour un électron de masse $m_e = 9,1 \cdot 10^{-31}$ kg, de charge $q = -1,6 \cdot 10^{-19}$ C, de vitesse $v_0 = 1,88 \cdot 10^7$ m.s⁻¹ et soumis au même champ magnétique :

$$R = \frac{9,1 \cdot 10^{-31} \times 1,88 \cdot 10^7}{1,6 \cdot 10^{-19} \times 0,1} = 1,1 \text{ mm}$$

La vitesse angulaire s'écrit :

$$\omega_c = \frac{qB}{m}$$

soit pour le proton :

$$\omega_c = \frac{1,6 \cdot 10^{-19} \times 0,1}{1,67 \cdot 10^{-27}} = 9,5 \cdot 10^6 \text{ rad.s}^{-1}$$

et pour l'électron :

$$\omega_c = \frac{1,6 \cdot 10^{-19} \times 0,1}{9,1 \cdot 10^{-31}} = 1,8 \cdot 10^{10} \text{ rad.s}^{-1}$$

L'électron tourne beaucoup plus rapidement que le proton et sur une trajectoire de rayon beaucoup plus petit.

4.4 Déviation d'une particule chargée par un champ magnétique

On se place dans le cas d'une trajectoire circulaire, c'est-à-dire dans le cas où la vitesse initiale $\vec{v}_0$ (qu'on prendra suivant Oy) est perpendiculaire au champ magnétique $\vec{B}$ qu'on prendra le long de l'axe Oz. On suppose que le champ magnétique n'existe que dans une zone limitée définie par $0 \leq y \leq L$. Expérimentalement il est difficile d'obtenir un champ magnétique uniforme sur une grande distance.

On cherche à déterminer l'angle de déviation θ de la trajectoire comme on l'a fait dans le cas d'un champ électrique.

$$\tan \theta = \tan(\widehat{\vec{v}_0, \vec{v}_1}) = \frac{\dot{x}_1}{\dot{y}_1}$$

La particule quitte la zone où règne le champ $\vec{B}$ en $y = L$. Or $y = \frac{v_0}{\omega_c} \sin \omega_c t$ (l'angle entre la vitesse initiale et le champ magnétique vaut $\frac{\pi}{2}$ et le sinus de cet angle 1). Donc la particule est en $y = L$ à l'instant :

$$t_1 = \frac{1}{\omega_c} \arcsin \frac{\omega_c L}{v_0}$$

Les composantes de la vitesse sont alors :

$$\begin{cases} \dot{x}_1 = v_0 \sin \omega_c t_1 \\ \dot{y}_1 = v_0 \cos \omega_c t_1 \end{cases}$$

Donc

$$\tan \theta = \frac{v_0 \sin \omega_c t_1}{v_0 \cos \omega_c t_1} = \tan \omega_c t_1$$

Soit

$$\theta = \omega_c t_1 = \arcsin \frac{\omega_c L}{v_0}$$

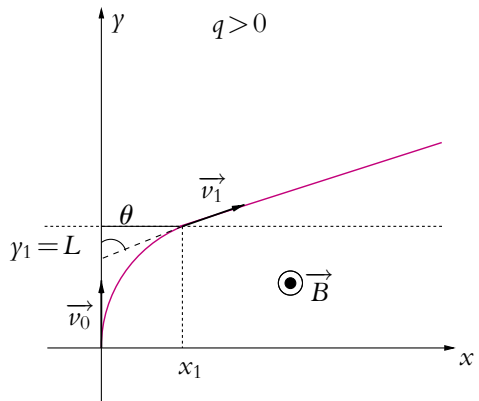


Figure 53.9 Déviation d'une particule chargée par un champ magnétique.

Application numérique : on cherche à déterminer la largeur L sur laquelle il faut appliquer le champ magnétique $B = 0,1$ T pour avoir la même déviation que lors de la déviation par un champ électrique soit $\theta = 1,71^\circ$. La vitesse initiale est toujours prise à $v_0 = 1,88 \cdot 10^7$ m.s⁻¹.

$$L = \frac{v_0}{\omega_c} \sin \theta = \frac{mv_0}{qB} \sin \theta = \frac{9,1 \cdot 10^{-31} \times 1,88 \cdot 10^7}{1,6 \cdot 10^{-19} \times 0,1} \sin(1,71^\circ) = 0,032 \text{ mm}$$

au lieu des deux centimètres avec un champ électrique.

Un champ magnétique permet donc une déviation plus forte qu’un champ électrique (les deux valeurs du champ sont importantes).

4.5 Application à la télévision et aux moniteurs d’ordinateur

Cette forte déviation par un champ magnétique avait une application pratique pour la réalisation d’une télévision ou d’un moniteur d’ordinateur aux dimensions raisonnables avant l’avènement des écrans plats. Le principe de fonctionnement est le même que celui de l’oscilloscope en remplaçant le champ électrique permettant la déviation des faisceaux d’électrons par un champ magnétique.

L’intérêt est de pouvoir dévier suffisamment le faisceau sans avoir une zone trop importante où le champ est appliqué ou sans être obligé d’éloigner beaucoup l’écran d’observation. On réduisait ainsi la profondeur des téléviseurs et des moniteurs d’ordinateurs, ce qui était loin d’être négligeable pour l’utilisateur.

Le champ magnétique est créé à l’aide de bobines : son intensité est donc proportionnelle au courant qui traverse ces dernières. Comme la déviation est proportionnelle au champ magnétique, on en déduit que la déviation est proportionnelle au courant circulant dans la bobine. Les variations d’intensité permettent de faire varier la déviation du faisceau sur l’écran et de visualiser le mouvement.

5. Mouvement dans un champ électrique et magnétique

L’étude se fait comme dans les paragraphes précédents en tenant compte des deux champs, électrique et magnétique. On ne procédera pas à la description du mouvement de manière systématique, on se limitera à l’analyse de quelques expériences intéressantes.

5.1 Expérience de Thomson

Cette expérience a permis à Sir J.J. Thomson de déterminer le rapport $\frac{q}{m}$ pour les électrons.

Elle consiste à accélérer un faisceau d’électrons initialement immobiles par un champ électrique $\vec{E}_0$ comme c’est le cas pour l’oscilloscope par exemple. Ce faisceau est

ensuite dévié par un champ électrique $\vec{E}$ appliqué perpendiculairement au mouvement des électrons du faisceau sur une longueur a . On mesure alors la déviation d induite par ce champ électrique sur un écran situé à une distance L . (Cf. figure 53.10).

L'application du seul champ électrique conduit à une déviation d'un angle θ tel que :

$$\tan \theta = \frac{-qEa}{mv^2} = \frac{d}{L}$$

en négligeant a devant L . Donc l'expression de la vitesse est :

$$v = \sqrt{\frac{-qEaL}{md}} = \sqrt{\frac{|q|EaL}{md}}$$

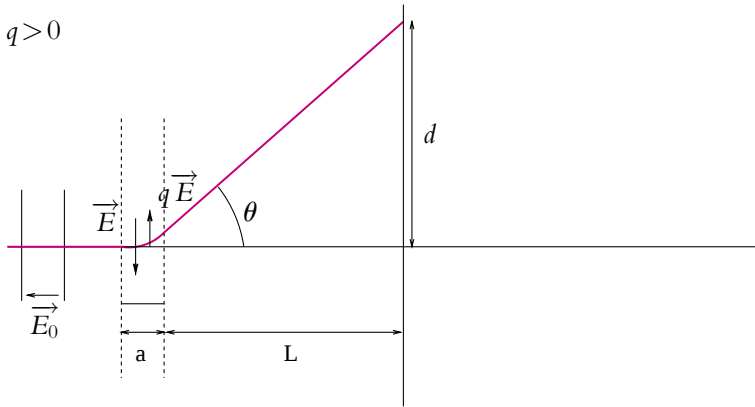


Figure 53.10 Expérience de Thomson avec le champ électrique seul.

On superpose alors un champ magnétique $\vec{B}$ perpendiculaire à $\vec{E}$ et tel que la force magnétique $q\vec{v} \wedge \vec{B}$ soit égale à la force électrique $q\vec{E}$. Le faisceau d'électrons n'est alors plus dévié.

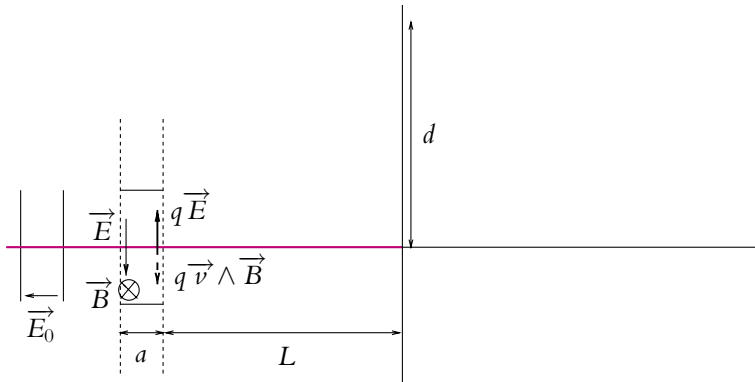


Figure 53.11 Expérience de Thomson avec champs électrique et magnétique.

L’égalité des forces électrique et magnétique donne : $qE = qvB$ soit $v = \frac{E}{B}$

En égalant les deux expressions de la vitesse, on obtient :

$$v^2 = \frac{E^2}{B^2} = \frac{-qEaL}{md}$$

ce qui permet d’obtenir le rapport :

$$\frac{|q|}{m} = \frac{dE}{aLB^2}$$

Cette expérience a permis de mesurer le rapport $\frac{|q|}{m}$ pour l’électron et conduit à de grands progrès dans la connaissance de l’électron et plus généralement de la constitution des atomes.

5.2 Spectromètre de masse

Le but d’un spectromètre de masse est d’analyser les ions présents dans un faisceau. On peut utiliser leurs différences de masse pour connaître leur proportion respective en supposant que leur charge est connue.

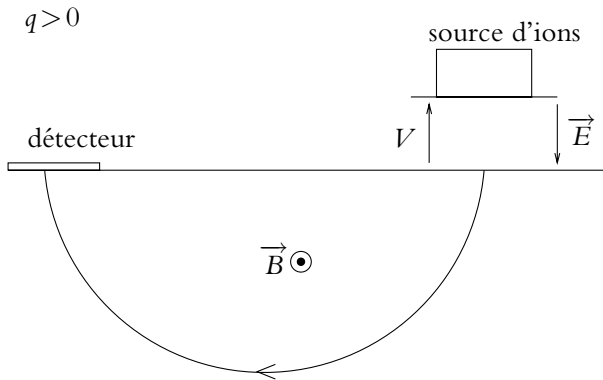


Figure 53.12 Spectromètre de masse.

Le montage utilisé comprend deux phases :

- Phase accélératrice :

On commence par accélérer les ions grâce à une forte différence de potentiel. L’application du théorème de l’énergie cinétique aux ions supposés initialement au repos fournit leur vitesse à la fin de cette phase :

$$v = \sqrt{\frac{2|q|V}{m}}$$

On se place cependant dans la limite classique en vérifiant que $v \ll c$ pour éviter un traitement relativiste.

- Phase de déviation :

On applique alors un champ magnétique $\vec{B}$ perpendiculairement au mouvement. Les ions décrivent alors une trajectoire circulaire de rayon R tel que :

$$R = \frac{mv}{|q|B}$$

en appliquant les relations établies précédemment.

On en déduit :

$$v = \frac{qBR}{m} = \sqrt{\frac{2|q|V}{m}}$$

soit :

$$\frac{q^2 B^2 R^2}{m^2} = \frac{2|q|V}{m} \quad \text{et} \quad \frac{|q|}{m} = \frac{2V}{B^2 R^2}$$

L'impact sur le détecteur permet de connaître le rayon de la trajectoire et donc le rapport $\frac{q}{m}$ connaissant la tension accélératrice et la norme du champ magnétique. Ce système permet de séparer des éléments ionisés en fonction de leur masse et de leur charge. Si on connaît par ailleurs la charge des ions, on peut en déduire leur masse. Ceci explique le nom de spectromètre de masse donné à ce dispositif.

Ordre de grandeur de la taille du dispositif :

La charge des ions sera de quelques charges élémentaires donc de l'ordre de 10^{-19} C. Leur masse molaire varie entre quelques centaines et quelques milliers de grammes par mole soit de l'ordre de 10^{-24} kg par ion. On prend une tension accélératrice de 10 kV et un champ magnétique de 0,1 T. On obtient alors un rayon de l'ordre de :

$$R = \frac{1}{B} \sqrt{\frac{2mV}{|q|}} = \frac{1}{0,1} \sqrt{\frac{210^{-24} \times 10^4}{10^{-19}}} = 4 \text{ m}$$

5.3 Cyclotron

L'existence d'une trajectoire circulaire des particules chargées soumises à un champ magnétique permet d'envisager des accélérateurs de particules agissant de manière cyclique. C'est le principe des cyclotrons dont le dispositif est le suivant :

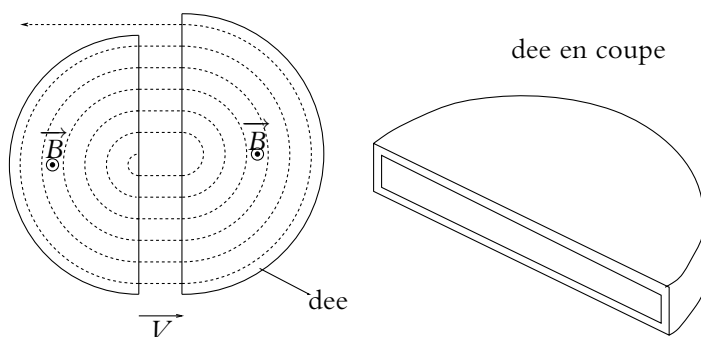


Figure 53.13 Constitution d'un cyclotron.

Ils sont constitués de deux demi-cylindres creux dénommés « *dees* » du fait de leur forme de « D » majuscules. Chaque dee est soumis à un champ magnétique uniforme, le sens du champ magnétique est le même dans les deux dees. Les deux dees sont séparés par un espace dans lequel on applique une différence de potentiel V . Souvent les particules injectées dans le dispositif ont été préalablement accélérées.

Dans chaque dee, la particule suit un mouvement circulaire lié au champ magnétique. Sa vitesse angulaire vaut :

$$\omega = \omega_c = \frac{qB}{m}$$

C’est bien la pulsation cyclotron dont on peut déduire la fréquence cyclotron :

$$f_c = \frac{\omega_c}{2\pi} = \frac{qB}{2\pi m}$$

Quand la particule traverse l’interstice entre les deux dees, elle est accélérée par la présence de la différence de potentiel V . Pour qu’elle soit effectivement accélérée à chaque demi-tour, il faut que les tensions aux instants correspondant à deux passages successifs dans l’interstice soient en opposition de phase. Sinon elle pourra être alternativement accélérée et décélérée. Pour obtenir une telle tension, il suffit de prendre une tension sinusoïdale dont la période est celle du cyclotron.

De cette manière, le rayon de la trajectoire augmente à chaque traversée de l’interstice puisqu’il est proportionnel à la vitesse.

On obtient le type de trajectoire indiquée en pointillés sur la figure précédente.

L’intérêt d’un tel dispositif est de pouvoir accélérer très fortement les particules sans être contraint d’appliquer une différence de potentiel énorme.

Pour le cyclotron de University of Michigan, Ann Arbor (près de Detroit), le champ magnétique vaut $B = 0,1 \text{ T}$ (on peut atteindre au maximum $1,5 \text{ T}$ mais on est alors dans le cadre relativiste), le diamètre des dees est $2,1 \text{ m}$.

Le problème essentiel des cyclotrons est d’obtenir un champ magnétique uniforme sur une surface aussi grande. Il s’agit d’un facteur limitant beaucoup sa taille et donc son utilisation.

On a alors une fréquence cyclotron pour les protons de :

$$f_c = \frac{qB}{2\pi m} = \frac{1,6 \cdot 10^{-19} \times 0,1}{2\pi \times 1,67 \cdot 10^{-27}} = 1,52 \text{ MHz}$$

La vitesse maximale est :

$$v_{\max} = R_{\max} \omega = 2\pi \times 1,05 \times 1,52 \cdot 10^6 = 10^7 \text{ m.s}^{-1}$$

Pour atteindre une telle vitesse à l'aide d'un seul champ électrique, il faudrait une différence de potentiel de :

$$V = \frac{mv^2}{2|q|} = \frac{1,67 \cdot 10^{-27} \times (10^7)^2}{2 \times 1,6 \cdot 10^{-19}} = 528 \text{ kV}$$

Le cyclotron n'est pas le seul type d'accélérateurs de particules. Dans tous les cas, l'accélération est due au champ électrique et le champ magnétique permet de faire revenir les particules au même point et donc de « refermer » les trajectoires.

6. Effet Hall

6.1 Étude du phénomène

Cet effet a été découvert par Edwin Herbert Hall (1855-1938) en 1880 : il a mis en évidence l'existence d'une tension perpendiculaire à un fil parcouru par un courant d'intensité I et soumis à l'action d'un champ magnétique perpendiculairement au fil. Il a également obtenu des mesures quantitatives du phénomène mettant en évidence la proportionnalité de la tension à I et à B et son caractère inversement proportionnel à la largeur du fil.

On étudie ici le modèle classique permettant d'expliquer cet effet.

On considère une plaque de métal supposée rectangulaire et d'épaisseur négligeable. Un courant d'intensité I la parcourt et on soumet l'ensemble à un champ magnétique perpendiculaire à la plaque.

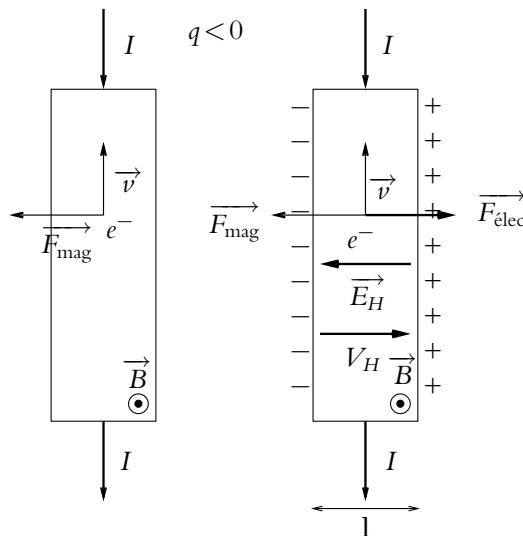


Figure 53.14 Effet Hall pour des charges négatives.

On suppose que les porteurs de charges sont des électrons : ils sont animés d’un mouvement rectiligne de vitesse $\vec{v}$ dans le sens opposé à I par définition du sens conventionnel du courant.

Quand on applique le champ magnétique, ils subissent une force d’origine magnétique :

$$\vec{F} = q\vec{v} \wedge \vec{B}$$

qui dévie leur trajectoire.

Les électrons ont donc tendance à s’accumuler sur un des côtés de la plaque de métal créant une accumulation de charges négatives de ce côté et une accumulation de charges positives de l’autre.

L’accumulation de charges de part et d’autre de la plaque provoque l’apparition d’un champ électrique $\vec{E}_H$ liée à ce déséquilibre.

Les électrons subissent alors l’action conjuguée de deux forces : l’une magnétique $\vec{F}_{\text{mag}} = q\vec{v} \wedge \vec{B}$ liée au champ magnétique $\vec{B}$, l’autre électrique $\vec{F}_{\text{elec}} = q\vec{E}_H$ due à l’apparition du champ électrique $\vec{E}_H$ par déséquilibre des charges.

Du fait des orientations relatives des deux champs, la force électrique est de même direction mais de sens opposé à la force magnétique. La force électrique tend donc à s’opposer à l’accumulation de charges.

Cependant tant que le module du champ électrique $\vec{E}_H$ est inférieur au module de $\vec{v} \wedge \vec{B}$, la force résultante est dans le sens du champ magnétique $\vec{B}$ et les électrons sont toujours déviés dans le même sens que précédemment. Cela augmente le déséquilibre et donc le module du champ $\vec{E}_H$. Il s’agit du régime transitoire.

Arrive en effet un instant où les normes de $\vec{E}_H$ et $\vec{v} \wedge \vec{B}$ deviennent égales. Les forces d’origine électrique et d’origine magnétique s’équilibrent et les deux effets magnétique et électrique (dû au déséquilibre des charges) se compensent. On a atteint le régime permanent : la force résultante est nulle et les électrons ne sont plus déviés.

On en déduit que :

$$q\vec{E}_H + q\vec{v} \wedge \vec{B} = \vec{0}$$

soit

$$\vec{E}_H = -\vec{v} \wedge \vec{B}$$

On a donc apparition d’une tension dite *tension de Hall*

$$V_H = E_H l$$

Si les porteurs sont des charges positives alors les orientations des différents vecteurs sont celles indiquées sur la figure 53.15.

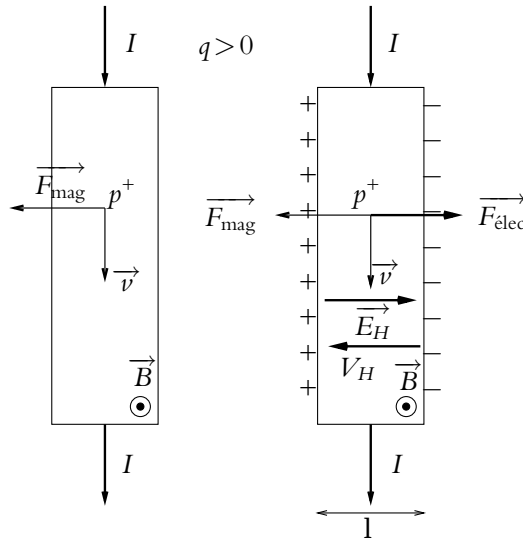


Figure 53.15 Effet Hall pour des charges positives.

Quel que soit le signe des porteurs de charges, on assiste à un déséquilibre des charges. Cependant on remarquera que l'accumulation de charges se fait de manière opposée à celle obtenue avec des électrons.

Le principe est le même puisque les forces d'origine magnétique et d'origine électrique ont la même orientation que la charge soit positive ou négative.

On a donc un régime transitoire pendant lequel les charges s'accumulent sur les deux faces et un régime permanent où la force de Hall compense la force magnétique et où les porteurs de charges ne sont plus déviés.

En notant n le nombre de porteur de charges, l la largeur de la plaque et ϵ son épaisseur, on obtient : $\vec{v} = \frac{1}{nq} \vec{j}$ et $j = \frac{I}{l\epsilon}$ soit :

$$V_H = \frac{1}{nq} \frac{I}{l\epsilon} Bl = R_H \frac{IB}{\epsilon}$$

en notant la constante de Hall :

$$R_H = \frac{1}{nq}$$

Cette constante de Hall est une grandeur algébrique dépendant du signe des porteurs de charges.

Ce modèle rend relativement bien compte de la réalité pour des métaux comme le cuivre, l'argent, l'or, le sodium ou le magnésium. Par exemple, pour le cuivre, on a pour les électrons $q = -e = -1,6 \cdot 10^{-19} \text{ C}$, $n = 8,5 \cdot 10^{28} \text{ m}^{-3}$ soit, d'après le modèle

précédent, une constante de Hall :

$$R_H = \frac{1}{8,5 \cdot 10^{28} \times (-1,6 \cdot 10^{-19})} = -0,7 \cdot 10^{-10} \text{ m}^3 \cdot \text{C}^{-1}$$

Expérimentalement, on mesure $R_H = -0,5 \cdot 10^{-10} \text{ m}^3 \cdot \text{C}^{-1}$, le modèle précédent donne le bon ordre de grandeur.

Par contre, pour d’autres métaux comme le fer ou le plomb, les valeurs expérimentales ne correspondent pas à celles du modèle notamment au niveau du signe. Pour expliquer ces différences, il faut recourir à la mécanique quantique, ce qui dépasse largement le cadre de ce cours.

6.2 Interprétation de la force de Laplace

L’effet Hall permet d’expliquer le mécanisme de transfert donnant lieu à la force de Laplace s’exerçant sur un conducteur parcouru par un courant I .

Le champ de Hall agit non seulement sur les porteurs de charges qui sont animés d’un mouvement mais également sur les charges fixes. L’action d’un champ magnétique sur des charges immobiles ne donne pas lieu à une force ; en revanche, un champ électrique induit une force sur toutes les charges qu’elles soient mobiles ou non : $\vec{F} = q\vec{E}$.

Par conséquent, les charges fixes du conducteur ne subissent pas d’action directe du champ magnétique mais l’existence de l’effet Hall implique qu’elles subissent la force :

$$\vec{F} = q\vec{E}_H = -q\vec{v} \wedge \vec{B}$$

Dans un élément de volume $d\tau$, on a une densité de charges fixes ρ_f :

$$\rho_f = \frac{dq}{d\tau} = -\rho_m$$

opposée à la densité de charges mobiles ρ_m du fait de la neutralité électrique du conducteur. Les charges mobiles subissent une force électrique et magnétique tandis que les charges fixes ne sont soumises qu’à la force électrique. La portion de conducteur correspondante subit donc une force :

$$d\vec{F} = \rho_f d\tau \vec{E}_H + \rho_m d\tau \vec{E}_H + \rho_m d\tau \vec{v} \wedge \vec{B} = \rho_m d\tau \vec{v} \wedge \vec{B} = \vec{j} \wedge \vec{B} d\tau$$

Ce qui se traduit pour un fil par :

$$d\vec{F} = I d\vec{l} \wedge \vec{B}$$

qui est l’expression de la force élémentaire de Laplace.

6.3 Application à la mesure de champ magnétique

On a établi que la tension de Hall vaut :

$$V_H = R_H \frac{IB}{e}$$

Elle est donc proportionnelle à la norme du champ magnétique. En amplifiant cette tension, on obtient un capteur capable de mesurer les champs magnétiques. Il s'agit des sondes à effet Hall qui pourront être utilisées en travaux pratiques.

Pour que la constante de Hall soit suffisamment grande pour rendre la mesure possible, on utilise en général un semi-conducteur à la place du métal (la densité volumique de porteurs n est beaucoup plus faible).

A. Applications directes du cours

1. Grandeurs cinétiques d'une particule chargée dans un champ magnétique

Une charge q de vitesse v pénètre dans une région où règne un champ magnétique constant et uniforme. Comparer les valeurs des quantités suivantes à l'entrée et à la sortie de la zone :

1. énergie cinétique,
2. quantité de mouvement,
3. moment cinétique par rapport au centre de la trajectoire.

2. Comparaison de l'action d'un champ magnétique sur un proton et sur un électron

Un électron et un proton de même énergie cinétique décrivent des trajectoires circulaires dans un champ magnétique uniforme. Comparer :

1. leur vitesse,
2. le rayon de leur trajectoire,
3. leur période.

3. Questions simples sur les particules dans champs $\vec{E}$ et $\vec{B}$, d'après concours A.T.S. 2004

On considère une particule ponctuelle, de charge q et de masse m , de vitesse initiale $\vec{V}_0$ à l'entrée d'une zone où règnent un champ électrique $\vec{E}$ ou un champ magnétique $\vec{B}$.

On suppose ces champs uniformes et indépendants du temps, et on néglige toute autre force que celles provoquées par ces champs.

1. a) La particule décrit une droite et possède une accélération constante a . Déterminer la direction et la norme du ou des champs qui provoquent cette trajectoire.
- b) Déterminer la position du point en fonction du temps.
2. a) La particule décrit une trajectoire circulaire de rayon R_0 , dans un plan xOy . Déterminer la direction du ou des champs qui provoquent cette trajectoire.
- b) Déterminer l'équation de la trajectoire et la relation entre la norme du champ, v_0 et R_0 . On suggère d'utiliser les coordonnées polaires.

4. Quelques mouvements de particules chargées dans $\vec{E}$ et $\vec{B}$

Dans tout l'exercice on se place dans un référentiel galiléen, associé à un repère cartésien $(O, \vec{u}_x, \vec{u}_y, \vec{u}_z)$. On étudie le mouvement d'une particule ponctuelle de masse m et de charge q dans les différents cas suivants sachant qu'on néglige l'effet de la pesanteur. On posera $\omega_c = \frac{qB}{m}$ (pulsation cyclotron). À l'instant initial, la particule est en O avec une vitesse initiale $\vec{v}_0$.

1. On suppose $\vec{E} = E\vec{u}_x$, $\vec{B} = B\vec{u}_x$ et $\vec{v}_0 = v_0\vec{u}_z$.
- a) Établir les équations du mouvement en projection sur les axes.

- b) Par intégrations, déterminer l'expression $x(t)$.
 - c) Avec les deux autres équations, déterminer une équation différentielle uniquement sur z . En déduire $z(t)$.
 - d) Déterminer $y(t)$.
 - e) Quelle est l'allure de la trajectoire ,
2. On suppose $\vec{E} = E\vec{u}_y$, $\vec{B} = B\vec{u}_z$ et $\vec{v}_0 = v_0\vec{u}_x$.
- a) Établir les équations du mouvement en projection sur les axes.
 - b) Montrer que le mouvement a lieu dans un plan.
 - c) On définit le nombre complexe $\underline{C} = x + iy$ avec $i^2 = -1$ pour résoudre les équations. Combiner les équations en x et y pour obtenir une équation différentielle en $\underline{C}$.
 - d) Résoudre cette équation pour déterminer $\underline{C}(t)$.
 - e) Par intégration déterminer $\underline{C}(t)$, puis $x(t)$ et $y(t)$.
 - f) Montrez que la trajectoire est rectiligne pour une certaine valeur de v_0 . (Ce dispositif sert à sélectionner les particules en fonction de leur vitesse).

5. Déflexion électrique dans un oscilloscope, d'après Banque PT 2000

Dans tout l'exercice on se place dans un référentiel galiléen, associé à un repère cartésien $(O, \vec{u}_x, \vec{u}_y, \vec{u}_z)$. Une zone de champ électrique uniforme (voir figure (53.16)) est établie entre les plaques P_1 et P_2 (le champ est supposé nul en dehors et on néglige les effets de bord) ; la distance entre les plaques est d , la longueur des plaques D et la différence de potentiel est $U = V_{P_1} - V_{P_2}$ positive. Des électrons (charge $q = -e$, masse m) accélérés pénètrent en O dans la zone de champ électrique uniforme avec une vitesse $\vec{v}_0 = v_0\vec{u}_z$ selon l'axe Oz .

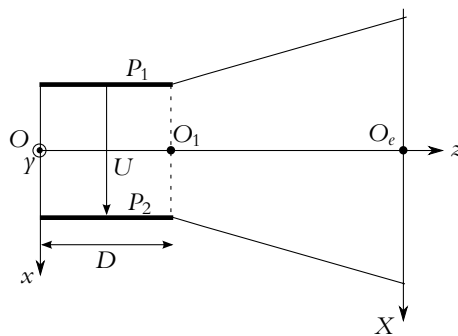


Figure 53.16

- 1. Calculer la force subie par les électrons en fonction de U , q , d et $\vec{u}_x$.
- 2. a) Déterminer l'expression de la trajectoire $x = f(z)$ de l'électron dans la zone du champ en fonction de d , U et v_0 .

- b) Déterminer le point de sortie K de la zone de champ ainsi que les composantes de la vitesse en ce point.
- c) Montrer que dans la zone en dehors des plaques, le mouvement est rectiligne uniforme.
- d) On note L la distance $O_1 O_e$ (figure 53.16). Déterminer l'abscisse X_P du point d'impact P de l'électron sur l'écran en fonction de U , v_0 , D , d et L .

B. Exercices et problèmes

1. Calculs de résistances

Calculer la résistance :

- 1. d'une couche cylindrique de conducteur ohmique de hauteur h comprise entre deux rayons R_1 et R_2 .
- 2. d'une couche sphérique de conducteur ohmique comprise entre deux rayons R_1 et R_2 ,

2. Conduction dans un métal en présence d'un champ magnétique

L'espace est repéré par les axes cartésiens Ox , Oy , Oz et les vecteurs $(\vec{u}_x, \vec{u}_y, \vec{u}_z)$. On s'intéresse à la conduction dans un métal assurée par des électrons de charge $q = -e$, de masse m , de densité n . Le métal est placé dans un champ électrique $\vec{E} = E_x \vec{u}_x + E_y \vec{u}_y + E_z \vec{u}_z$ et un champ magnétique $\vec{B} = B \vec{u}_z$. Ces champs sont uniformes et permanents. L'interaction entre les électrons et les ions fixes du réseau est modélisée par une force de frottement : $\vec{F}_v = -\frac{m\vec{v}}{\tau}$, $\vec{v}$ étant la vitesse des électrons.

- 1. Déterminer l'unité de τ .
- 2. a) Établir l'équation différentielle du mouvement d'un électron. Que devient cette équation en régime permanent ?
- b) On note $\vec{j} = nq\vec{v}$ le vecteur densité de courant. Montrer qu'en régime permanent, les composantes de ce vecteur vérifie le système d'équations :

$$\begin{cases} j_x - \frac{q\tau B}{m} j_y = \gamma E_x \\ \frac{q\tau B}{m} j_x + j_y = \gamma E_y \\ j_z = \gamma E_z \end{cases}$$

où γ est une constante que l'on déterminera.

- c) En déduire que l'on peut écrire j_x et j_y sous la forme suivante :

$$j_x = \gamma(aE_x + bE_y) \quad \text{et} \quad j_y = \gamma(-bE_x + aE_y)$$

et où a et b sont deux constantes à déterminer en fonction des données.

3. Dans le cas dans lequel la conduction ne peut avoir lieu que suivant Ox , montrer que la présence du champ magnétique $\vec{B}$ impose la présence d'un champ électrique (appelé champ de Hall) et déterminer sa direction.

Calculer la constante de Hall : $\frac{E_y}{j_x B}$.

3. Conductivité complexe en régime variable

On considère un métal dans lequel la conduction est assurée par des électrons (masse m , charge $q < 0$, densité n) soumis à un champ électrique $\vec{E} = E_0 \cos \omega t \vec{u}_x$ ($E_0 > 0$). On modélise la conduction par le modèle de Drude : la gravité est négligée et l'interaction des électrons avec les ions fixes du réseau est modélisée par une force de frottement : $\vec{F}_v = -\frac{m\vec{v}}{\tau}$, $\vec{v}$ étant la vitesse des électrons et $\tau = 10^{-14}$ s.

1. Établir l'équation du mouvement d'un électron.

2. On définit le champ électrique complexe $\vec{E} = E_0 \exp(i\omega t) \vec{u}_x$ avec $i^2 = -1$. En adoptant en régime forcé de pulsation ω une représentation complexe de la vitesse : $\vec{v} = \vec{v}_0 \exp(i(\omega t - \varphi))$, déterminer les expressions de $\vec{v}$, $\vec{v}_0$ et $\tan \varphi$ (on rappelle que $q < 0$).

3. En déduire la représentation complexe du vecteur densité de courant $\vec{j} = nq\vec{v}$, définir la conductivité complexe $\underline{\gamma}$ telle que $\vec{j} = \underline{\gamma}\vec{E}$ et l'exprimer en fonction de la conductivité en continu γ_0 .

4. Interpréter le résultat en ce qui concerne la validité de la loi d'Ohm en régime non permanent.

4. Modèle de Thomson, d'après Centrale MP 2005

En 1904 le physicien J.J. Thomson propose le modèle suivant pour l'atome d'hydrogène :

- Il est constitué d'une sphère de centre O et de rayon a ;
- la charge positive e de l'atome est répartie dans le volume intérieur de cette sphère ;
- la sphère est supposée fixe dans un référentiel galiléen, auquel on associe le repère $(O, \vec{u}_x, \vec{u}_y, \vec{u}_z)$;
- on repère par $\vec{r} = \overrightarrow{OM}$ la position de l'électron (masse m_e) qui se déplace à l'intérieur de la sphère ;
- on néglige l'interaction gravitationnelle.

1. Quelle est la force ressentie par l'électron ? On posera $k = \frac{e^2}{4\pi\epsilon_0 a^3}$

2. a) Montrer que le mouvement de l'électron est plan.

b) La loi de force précédente définit un modèle, analogue à trois dimensions de l'oscillateur harmonique. Déterminer les équations du mouvement de l'électron, en projection sur les axes.

c) Déterminer la trajectoire de l'électron pour les conditions initiales suivantes : à $t = 0$, $\vec{r}_0 = r_0 \vec{u}_x$ et $\vec{v}_0 = v_0 \vec{u}_z$

d) Quelle est la période du mouvement en fonction de m_e et k ?

5. Mouvement de gouttelettes chargées, d'après ENAC 2005

On disperse un brouillard de fines gouttelettes sphériques d'huile, de masse volumique $\rho_h = 1,3 \cdot 10^3 \text{ kg.m}^{-3}$, dans l'espace séparant les deux plaques horizontales d'un condensateur plan, distantes de $d = 2 \cdot 10^{-2} \text{ m}$. Les gouttelettes sont chargées négativement et sans vitesse initiale.

Toutes les gouttelettes ont même rayon R mais pas forcément la même charge $-q$. En l'absence de champ électrique E , une gouttelette est soumise à son poids ($g = 9,81 \text{ m.s}^{-2}$), à la poussée d'Archimède de l'air ambiant de masse volumique $\rho_a = 1,3 \text{ kg.m}^{-3}$ et à une force de frottement visqueux $\vec{f} = -k\vec{v}$, avec $k = \alpha R$ et $\alpha = 3,4 \cdot 10^{-4} \text{ S.I.}$

L'accélération de la pesanteur $\vec{g}$ sera prise égale à $9,81 \text{ m.s}^{-2}$.

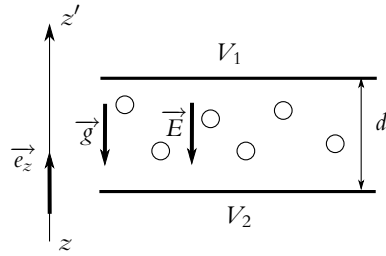


Figure 53.17

1. a) Déterminer la vitesse limite $\vec{v}_0$
- b) Déterminer l'expression de la vitesse des gouttes $\vec{v}(t)$. On fera apparaître un temps caractéristique τ .
- c) On mesure $v_0 = 2 \cdot 10^{-4} \text{ m.s}^{-1}$, déterminer la valeur de k .
2. On applique une différence de potentiel $U = V_1 - V_2$ de manière à avoir un champ électrique $\vec{E}$ dirigé vers le bas
- a) Déterminer E en fonction de U
- b) Une gouttelette est immobilisée pour $U = 3200 \text{ V}$. Calculer la charge q .

6. Mouvement d'une particule chargée dans un champ magnétique uniforme et non uniforme, d'après Centrale TSI 2005

La première partie est plutôt facile, la deuxième partie est difficile.

Une particule de masse m et de charge q , est soumise à l'action d'un champ magnétique $\vec{B}$. Les axes Ox, Oy, Oz sont repérés par les vecteurs $\vec{u}_x, \vec{u}_y, \vec{u}_z$.

1. Dans cette question, le champ $\vec{B} = B\vec{u}_z$ est uniforme et permanent, avec $B > 0$. On pose $\omega = \frac{qB}{m}$. La vitesse de la particule est $\vec{v} = v_x\vec{u}_x + v_y\vec{u}_y + v_z\vec{u}_z$. Ainsi, $\vec{v}_L = v_L\vec{u}_z$ et $\vec{v}_\perp = v_x\vec{u}_x + v_y\vec{u}_y$ désignent respectivement les composantes parallèle et perpendiculaire au champ. On note $v_\perp = \|\vec{v}_\perp\|$. À l'instant initial, la particule se trouve en O avec la vitesse $\vec{v}_0 = v_{\perp 0}\vec{u}_x + v_{L0}\vec{u}_z$ avec $v_{\perp 0} > 0$ et $v_{L0} > 0$.

- a) Montrer que l'énergie cinétique E_c est constante au cours du temps.
- b) Montrer que $\vec{v}_L$ est une constante du mouvement. En déduire que $v_\perp$ est également constante au cours du mouvement. On pose $E_{c\perp} = \frac{1}{2}mv_\perp^2$.

On étudie la projection du mouvement de la particule dans le plan $P_\perp$ perpendiculaire à $\vec{B}$.

- c) Déterminer les composantes v_x et v_y de la particule en fonction de $v_{\perp 0}$, ω et du temps t .

d) Montrer que la projection de la trajectoire de la particule dans le plan $P_\perp$ est un cercle Γ . Déterminer le centre C (centre guide) et le rayon a (rayon de giration) du cercle, ainsi que la période de révolution T_1 de la particule sur ce cercle en fonction de $v_{\perp 0}$ et ω .

e) L'orbite circulaire Γ peut être assimilée à une petite spire de courant ; déterminer l'intensité i de ce courant associé au mouvement de la particule. En déduire le moment dipolaire magnétique $\vec{\mu} = -Si\vec{u}_z = -\mu\vec{u}_z$ avec $\mu > 0$.

Exprimer μ en fonction de $E_{c\perp}$ et B puis en fonction de q , m et du flux φ du champ $\vec{B}$ à travers l'orbite circulaire Γ .

f) Quelle est la trajectoire de la particule chargée ? Expliquer pourquoi elle s'enroule sur un tube de champ du champ $\vec{B}$.

On peut décomposer le mouvement de la particule en un mouvement sur un cercle dont le centre C (centre guide) se déplace à la vitesse $\vec{v}_L$ le long de Oz . Quelle distance b parcourt le centre C sur Oz durant la période T_1 ? Exprimer b en fonction de v_L et ω .

2. On suppose que le champ $\vec{B}$ n'est plus tout à fait uniforme, ses variations restent très faibles sur une distance de l'ordre du rayon de giration a ou de la distance b . Le champ $\vec{B}$ présente la symétrie de révolution d'axe Oz et on admet que la composante B_z ne dépend que de z et est positive, dans la zone située au voisinage de l'axe Oz où se déplace la particule chargée et dans laquelle il n'y a aucun courant. Un point M est repéré par ses coordonnées cylindriques : $M(r, \theta, z)$.

a) Montrer que la composante orthoradiale B_θ du champ $\vec{B}$ est nulle.

b) En exprimant la flux du champ $\vec{B}$ à travers un petit cylindre judicieusement choisi, déterminer la composante radiale B_r au point M en fonction de r et de la dérivée $\frac{dB_z}{dz}$.

En considérant la champ $\vec{B}$ localement uniforme on peut utiliser certains résultats de la partie (1) :

- une particule chargée décrit un mouvement circulaire dans un plan perpendiculaire à $\vec{B}$, autour d'un centre guide C se déplaçant autour de Oz .
- En raison de la variation de $\vec{B}$ le rayon a du cercle et la période T_1 varient.
- Les expressions de a et T_1 trouvées à la question (1) restent valables en remplaçant B par B_z .
- Le déplacement de C selon Oz n'est plus uniforme (v_L et $v_\perp$ ne sont plus constantes).

c) Montrer que la composante F_z sur l'axe Oz de la force qui agit sur la particule chargée a pour expression :

$$F_z = -\frac{mv_\perp^2}{2B_z} \frac{dB_z}{dz}$$

d) On peut admettre le résultat de cette question pour la suite.

En utilisant l'expression de μ obtenue à la question (1.e) dans laquelle B est remplacée par B_z montrer que :

$$\frac{d\mu}{dt} = \frac{1}{B_z} \frac{dE_{c\perp}}{dt} - \frac{E_{c\perp}}{B_z^2} \frac{dB_z}{dz} \dot{z} = \frac{1}{B_z} \frac{dE_c}{dt}$$

En déduire que μ est une constante du mouvement. Peut-on encore dire que la particule s'enroule sur un tube du champ $\vec{B}$ au cours de son mouvement ?

- e) Exprimer l'énergie cinétique E_c de la particule en fonction de m , v_L , μ et B_z .
- f) En déduire que la particule chargée ne peut entrer dans une zone où la composante B_z du champ dépasse une valeur maximale B_{max} que l'on exprimera en fonction de E_c et μ .
- g) Sur la figure (53.18), ont été représentées plusieurs lignes de champ dans le plan méridien passant par Oz . On constate que le plan xOy est un plan de symétrie pour la distribution de courants créant le champ $\vec{B}$. Comment varie l'intensité de $\vec{B}$ de O à M_1 , et de O à M'_1 ? Que peut-on dire de l'intensité du champ B_O en O ? Quelle relation y-a-t-il entre les intensités des champs en M_1 et M'_1 qui sont symétriques par rapport à O ?
- h) On suppose B_{max} compris entre B_O et B_1 (champ en M_1). Montrer que dans ces conditions, le centre guide C oscille périodiquement entre deux points M_0 et M'_0 , symétriques par rapport au point O , d'abscisses respectives $z_0 > 0$ et $-z_0$, et que la période T_2 de ce mouvement est donnée par :

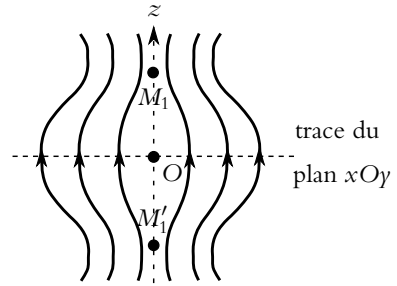


Figure 53.18

$$T_2 = 4\sqrt{\frac{m}{2\mu}} \int_0^{z_0} \frac{dz}{\sqrt{B_{max} - B(z)}}$$

où $B(z)$ représente l'intensité du champ magnétique en un point M variable de l'axe Oz situé entre O et M_0 .

7. Étude de l'effet Hall en régime indépendant du temps, d'après CCP PC 2002

On considère une plaque rectangulaire d'épaisseur h et de largeur b . Elle est réalisée dans un semi-conducteur de type N où la conduction électrique est assurée par des électrons mobiles dont le nombre par unité de volume est n . On notera par e la charge élémentaire égale à $1,6 \cdot 10^{-19}$ C. La plaque est parcourue par un courant électrique d'intensité I , uniformément réparti sur la section de la plaque avec la densité volumique $\vec{J} = J\vec{u}_x$ ($J > 0$) (fig. 1). Elle est alors placée dans un champ magnétique uniforme $\vec{B} = B\vec{u}_z$ avec $B > 0$ créé par des sources extérieures. Le champ magnétique créé par le courant dans la plaque est négligeable devant $\vec{B}$.

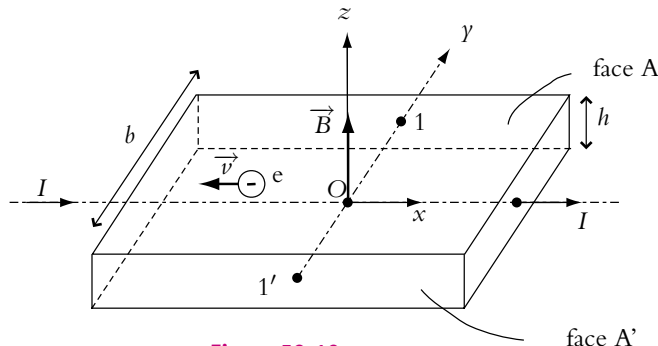


Figure 53.19

On suppose qu'en présence du champ magnétique $\vec{B}$, le vecteur densité de courant est toujours égal à $\vec{J} = J\vec{u}_x$.

1. Exprimer le vecteur vitesse $\vec{V}$ des électrons dans la plaque en fonction de $\vec{J}$, n et e . Montrer qu'en présence du champ magnétique $\vec{B}$ en régime permanent, il apparaît un champ électrique appelé champ électrique de Hall : $\vec{E}_H = \frac{1}{ne} \vec{J} \wedge \vec{B}$. Exprimer les composantes de $\vec{E}_H$.

2. On considère deux points 1 et 1' en vis-à-vis des faces A et A' de la plaque. Calculer la différence de potentiel $U_H = V(1) - V(1')$ appelée tension de Hall. Montrer que U_H peut s'écrire :

$$U_H = \frac{C_H}{h} IB$$

Expliciter la constante C_H .

A.N. : Pour l'antimoniure d'indium, $C_H = 375.10^{-6} \text{ m}^3 \cdot \text{C}^{-1}$, de plus : $I = 0,1 \text{ A}$, $h = 0,3 \text{ mm}$, $B = 1 \text{ T}$.

Calculer U_H ainsi que la densité volumique n , en électrons par mètre cube.

3. On veut établir la loi d'Ohm locale, c'est-à-dire, la relation entre le champ électrique $\vec{E}$ dans la plaque et la densité du courant $\vec{J}$ en présence du champ magnétique $\vec{B}$. Soit $\vec{E}' = E'\vec{u}_x$ la partie du champ électrique colinéaire à $\vec{J}$. On pose $\vec{J} = \sigma \vec{E}'$, σ étant une grandeur positive.

Quelle caractéristique du matériau de la plaque σ représente-t-elle ?

Montrer qu'en présence du champ magnétique, on a $\vec{J} = \sigma(\vec{E} - C_H \vec{J} \wedge \vec{B})$.

4. Tracer dans un plan xOy de la plaque les vecteurs $\frac{\vec{J}}{\sigma}$, $\vec{E}$ et $C_H \vec{J} \wedge \vec{B}$ et les lignes équipotentiellles en présence puis en absence de champ magnétique. Faire deux figures en vue de dessus par rapport à la figure 1.

5. Soit θ l'angle entre les vecteurs $\vec{J}$ et $\vec{E}$. Montrer que l'angle θ ne dépend que de B et du semi-conducteur. Préciser le domaine de définition de θ pour le semi-conducteur étudié.

6. On veut utiliser la plaque pour mesurer l'induction magnétique B , en mesurant la tension de Hall U_H . Il faut donc qu'en absence du champ magnétique $U_H(B=0) = U_{H0} = 0$. Pour cela, il faut souder deux fils conducteurs exactement en vis-à-vis. C'est un problème difficile, vu les dimensions de la plaque. Proposer un schéma de montage, utilisant un potentiomètre, ainsi que le protocole expérimental qui permet d'avoir $U_{H0} = 0$.

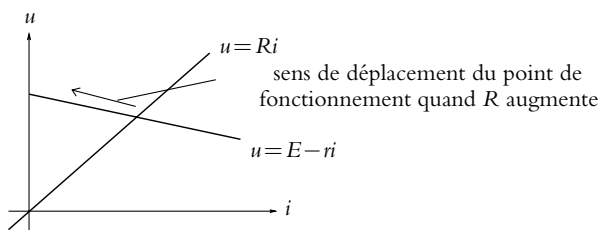
Solutions des exercices

Première période

Partie I – Électrocinétique I

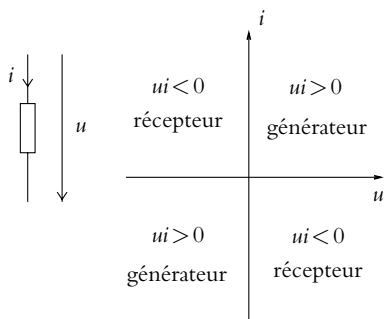
Chapitre 2

- A.1** 1. La source de tension étant idéale, la tension reste constante quoi qu'il arrive.
2. L'appareil ne change pas de caractéristique donc le fait que la tension à ses bornes soit constante implique que le courant qui le traverse reste constant aussi.
3. On raisonne graphiquement en traçant pour chaque dipôle la courbe donnant la tension à ses bornes en fonction de l'intensité qui le traverse. Les valeurs de l'intensité et de la tension doivent vérifier les deux relations $u = Ri$ et $u = E - ri$ donc graphiquement on peut les obtenir par l'intersection des deux courbes. Le point d'intersection est appelé point de fonctionnement.

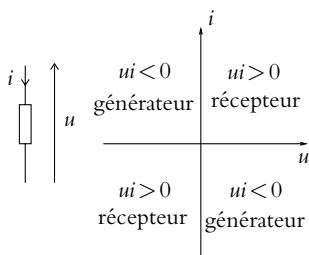


Quand R augmente, le point de fonctionnement se déplace dans le sens indiqué sur la figure précédente donc l'intensité du courant dans le circuit diminue et la tension aux bornes du générateur idéal également tandis que celle aux bornes de R augmente.

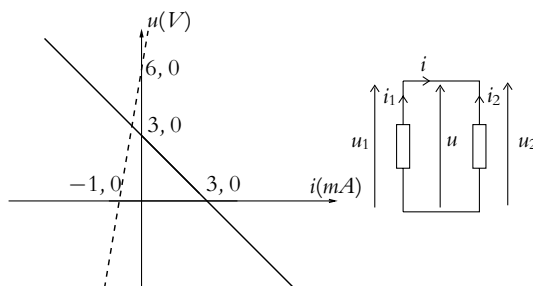
- A.2** 1. En convention générateur, la puissance fournie est donnée par $P_{\text{fournie}} = ui$. On a un générateur si la puissance fournie est effectivement fournie à savoir si $ui > 0$; sinon on a un récepteur. On en déduit donc :



2. De même, en convention récepteur, la puissance reçue est donnée par $P_{\text{reçue}} = ui$. On a un récepteur si la puissance reçue est effectivement reçue à savoir si $ui > 0$; sinon on a un générateur. On en déduit donc :



3. On effectue le choix d'une convention pour un dipôle. Dans la suite, on suppose par exemple qu'on prend la convention générateur pour le premier dipôle. Les sens de u et de i sont imposés alors pour le second dipôle se retrouvent dans l'exemple considéré en convention récepteur. Il n'est donc pas possible de prendre une unique convention pour tous les dipôles.
4. On trace la caractéristique du dipôle 1, ce qui ne pose pas de problème puisque sa caractéristique est fournie dans la « bonne » convention. En revanche, pour le dipôle 2, on a $i_2 = -i$ compte tenu de la convention prise. Il faut donc effectuer une symétrie de la caractéristique par rapport à l'axe des ordonnées :

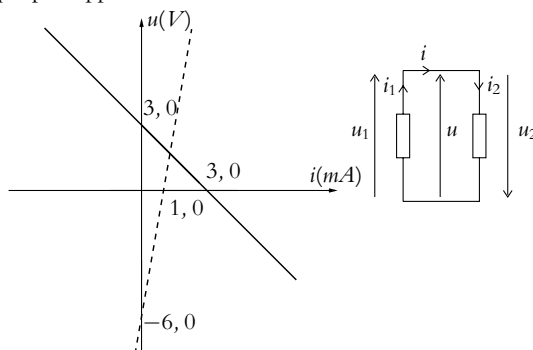


- a) Il convient de déterminer l'intersection des deux caractéristiques soit graphiquement c'est-à-dire en relevant les coordonnées du point d'intersection des deux droites soit en résolvant le système

$$\begin{cases} u = 3,0 - 1,0 \cdot 10^3 i \\ u = 6,0 + 6,0 \cdot 10^3 i \end{cases}$$

On obtient dans les deux cas $i = -0,43 \text{ mA}$ et $u = 3,4 \text{ V}$.

- b) On a la convention générateur pour le dipôle 1 et on obtient $ui < 0$ donc le dipôle 1 est récepteur. On a la convention récepteur pour le dipôle 2 et on obtient $ui < 0$ donc le dipôle 2 est générateur.
5. a) Il n'y a toujours pas de soucis pour le dipôle 1 dont la caractéristique est fournie dans la « bonne » convention mais pour le dipôle 2, on a $u_2 = -u$ et $i_2 = i$: il faut donc procéder à une symétrie de la caractéristique par rapport à l'axe des abscisses :



b) On procède ensuite comme à la question précédente en résolvant

$$\begin{cases} u = 3,0 - 1,0 \cdot 10^3 i \\ u = -6,0 + 6,0 \cdot 10^3 i \end{cases}$$

On obtient $i = 1,3 \text{ mA}$ et $u = 1,7 \text{ V}$.

c) On a la convention générateur pour le dipôle 1 et $ui > 0$ donc il est générateur.
On a la convention récepteur pour le dipôle 2 et $ui > 0$ donc il est récepteur.

A.3 1. Détermination de la résistance entre A et B.

La droite AB est un axe de symétrie donc les points C, D, E et F sont au même potentiel. Par conséquent, on peut relier ces points sans changer le comportement du circuit. Cela revient à court-circuiter les résistances CD, DE, EF et FC . On obtient le circuit de la figure (S2.1).

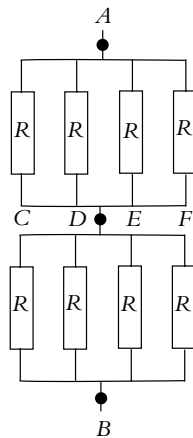


Figure S2.1

On a alors une résistance $R/4$ équivalente aux 4 résistances en parallèle entre A et le point central ; de même entre le point central et B . Finalement le circuit est équivalent à deux résistances $R/4$ en série soit une résistance totale égale à $R/2$.

2. Détermination de la résistance entre C et D.

Le plan $CDEF$ est plan de symétrie, donc A et B sont au même potentiel. Par conséquent, on peut les relier par un fil sans rien changer au circuit ; les résistors entre le point A et les points $CDEF$ se retrouvent en parallèle sur les résistors entre le point B et les points $CDEF$ ce qui donne le schéma équivalent de la figure (S2.2) où toutes les résistances sont égales à R . Par association des dipôles en parallèle, on obtient le circuit de la figure (S2.3).

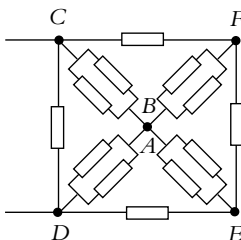


Figure S2.2

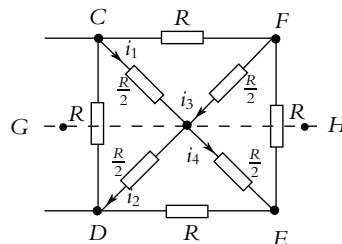


Figure S2.3

On peut alors considérer que le plan GH est un plan d'antisymétrie donc les intensités de la figure (S2.3) sont telles que $i_1 = i_2$ et $i_3 = i_4$. Ainsi le nœud AB peut être supprimé ce qui donne le schéma équivalent de la figure (S2.4). Par association des résistors en parallèle, on trouve le schéma de la figure (S2.5) ce qui donne un résistor de résistance $R/2$ en parallèle avec trois résistors R , R et $R/2$ en série, soit $R/2$ et $5R/2$ en parallèle, soit finalement une résistance $5R/12$.

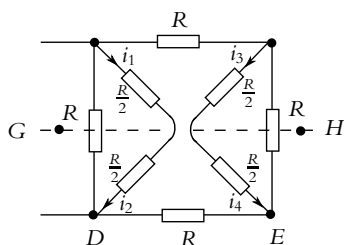


Figure S2.4

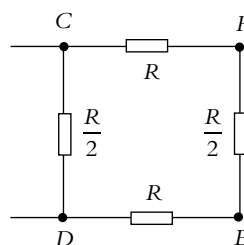


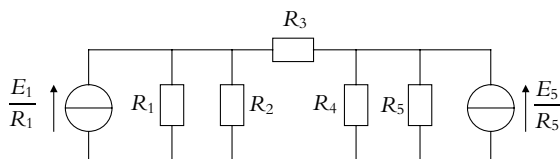
Figure S2.5

A.4 Les trois résistances de droite sur la figure sont en série donc équivalentes à $R_1 = 2R + R'$. Cette résistance R_1 est elle-même en parallèle sur la quatrième résistance R , ce qui donne une résistance équivalente $R_{eq} = \frac{R_1 R}{R + R_1}$. On veut que $R_{eq} = R'$ soit :

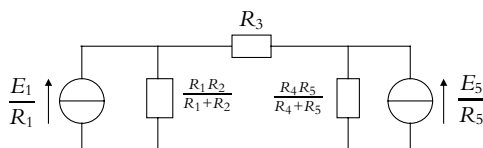
$$\frac{R_1 R}{R + R_1} = R' \Rightarrow \frac{R(2R + R')}{3R + R'} = R' \Rightarrow R'^2 + 2RR' - 2R^2 = 0$$

Les solutions pour R' de cette équation sont $R(-1 \pm \sqrt{3})$. La solution physique est la solution positive donc $R' = R(\sqrt{3} - 1)$.

A.5 1. On transforme tous les modèles de Thévenin en modèles de Norton :



On associe ensuite les résistances R_1 et R_2 ainsi que R_4 et R_5 qui sont en parallèle soit :



On revient aux modèles de Thévenin et on associe les modèles de Thévenin en série. On obtient un diviseur de tension avec R_3 et finalement

$$i = \frac{R_2 (R_4 + R_5) E_1 - R_4 (R_1 + R_2) E_5}{R_1 R_2 (R_4 + R_5) + R_3 (R_1 + R_2) (R_4 + R_5) + R_4 R_5 (R_1 + R_2)}$$

2. a) On remplace la source idéale de tension de force électromotrice E_5 par un fil et on applique la même méthode qu'à la première question. On obtient

$$i_1 = \frac{R_2 (R_4 + R_5) E_1}{R_1 R_2 (R_4 + R_5) + R_3 (R_1 + R_2) (R_4 + R_5) + R_4 R_5 (R_1 + R_2)}$$

- b) On remplace les indices 1, 2, 3, 4 et 5 par respectivement 5, 4, 3, 2 et 1 et on inverse le sens du courant dans la résistance R_3 . On en déduit

$$i_2 = - \frac{R_4 (R_1 + R_2) E_5}{R_1 R_2 (R_4 + R_5) + R_3 (R_1 + R_2) (R_4 + R_5) + R_4 R_5 (R_1 + R_2)}$$

- c) Le courant est alors par le principe de superposition la somme des courants i_1 et i_2 et on retrouve le résultat de la première question.

A.6 On ne modifie pas la branche contenant la résistance R et on transforme le générateur de Norton en un générateur de Thévenin de résistance R_0 et de force électromotrice $R_0 I_0$ pour pouvoir l'associer à l'autre générateur de Thévenin en série. On obtient le générateur de Thévenin équivalent de résistance $R_0 + R_2$ et de force électromotrice $E_2 + R_0 I_0$. On transforme alors les deux générateurs de Thévenin en générateurs de Norton de résistances respectives $R_0 + R_2$ et R_1 et de courant à vide respectifs $\frac{R_0 I_0 + E_2}{R_0 + R_2}$ et $\frac{E_1}{R_1}$. On associe alors ces deux générateurs de Norton en parallèle pour obtenir le générateur de Norton

équivalent de résistance $\frac{R_1 (R_0 + R_2)}{R_0 + R_1 + R_2}$ et de courant à vide $\frac{E_1}{R_1} - \frac{R_0 I_0 + E_2}{R_0 + R_2}$. On applique la relation du pont diviseur de courant pour en déduire l'intensité cherchée $i = \frac{(R_0 + R_2) E_1 - R_1 (R_0 I_0 + E_2)}{R (R_0 + R_1 + R_2) + R_1 (R_0 + R_2)}$.

Une autre méthode consiste à transformer le générateur de Norton en générateur de Thévenin avant d'appliquer la relation du pont diviseur de tension.

- A.7** 1. On commence par associer les deux résistances R de gauche en série avec la source de tension parfaite e_1 , ce qui donne une résistance $2R$ en série avec la même source de tension parfaite e_1 . On se retrouve avec des modèles de Thévenin $(e_1, 2R)$ et (e_2, R) en parallèle. On les transforme donc en modèles de Norton. On rappelle les résultats de la transformation Thévenin-Norton sur la figure (S2.6), ce qui donne le résultat de la figure (S2.7) pour le circuit.

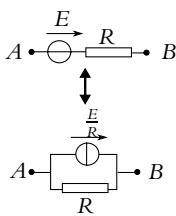


Figure S2.6

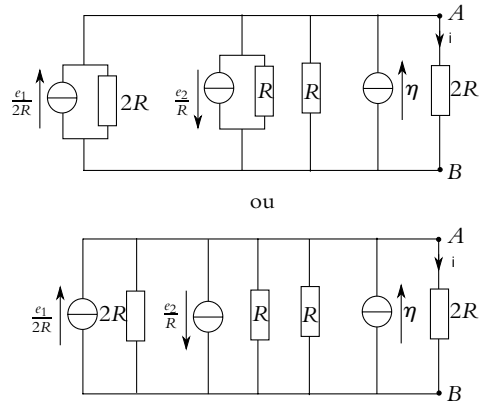


Figure S2.7

Finalement on a trois sources de courant parfaites en parallèle, équivalentes à une seule source de courant de court-circuit $I_N = \frac{e_1}{2R} - \frac{e_2}{R} + \eta$ et trois résistors en parallèle équivalents à R_N telle que

$$\frac{1}{R_N} = \frac{1}{2R} + \frac{1}{R} + \frac{1}{R} \Rightarrow R_N = \frac{2R}{5}$$

En remplaçant la partie gauche du circuit par son modèle de Norton équivalent, on obtient le circuit de la figure (S2.8).

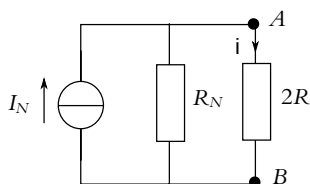


Figure S2.8

On calcule alors l'intensité i par un diviseur de courant en notant $G_N = 1/R_N$ et $G_2 = 1/(2R)$:

$$i = \frac{G_2}{G_2 + G_N} I_N = \frac{1}{6} \left(\frac{e_1}{2R} - \frac{e_2}{R} + \eta \right)$$

2. On utilise le modèle de Norton déterminé à la question précédente et on le transforme en modèle de Thevenin, ce qui donne le résultat de la figure (S2.9) avec $e_T = R_N I_N = (e_1 - 2e_2 + 2\eta R)/5$.

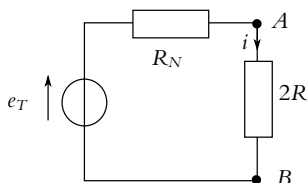


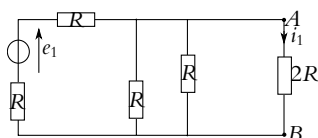
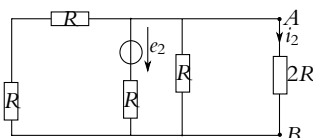
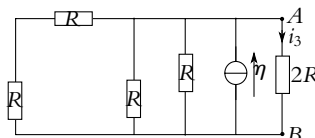
Figure S2.9

Une loi des mailles donne :

$$e_T = \left(\frac{2R}{5} + 2R \right) i = \frac{12R}{5} i$$

et l'on retrouve pour i le résultat précédent.

3. Pour le théorème de superposition, on éteint les générateurs les uns après les autres (on en garde un seul allumé à chaque état), sachant qu'une source de tension éteinte est équivalente à un court-circuit et qu'une source de courant éteinte est équivalente à un circuit ouvert. Les trois états obtenus sont représentés sur les figures (S2.11, S2.10 et S2.12).

Figure S2.10
 e_1 alluméFigure S2.11
 e_2 alluméFigure S2.12
 η allumé

Avec les associations de résistances, les trois circuits précédents se simplifient comme sur les figures (S2.13, S2.14 et S2.15) :

- Pour le premier circuit, les deux résistors en série de résistance R sont équivalents à un résistor de résistance $2R$. Les deux résistors en parallèle de résistance R sont équivalents à un résistor de résistance $R/2$.
- Pour le deuxième circuit, les deux résistors en série de résistance R sont équivalents à un résistor de résistance $2R$. On se retrouve alors avec un résistor de résistance $2R$ en parallèle avec un résistor de résistance R ce qui est équivalent à un résistor de résistance $\frac{2R \cdot R}{2R + R} = 2R/3$.
- Pour le troisième circuit, les deux résistors en série de résistance R sont équivalents à un résistor de résistance $2R$. On se retrouve alors avec un résistor de résistance $2R$ en parallèle avec deux résistors de résistance R ce qui est équivalent à un résistor de résistance R_{eq} telle que $\frac{1}{R_{eq}} = \frac{1}{2R} + \frac{1}{R} + \frac{1}{R}$ soit $R_{eq} = 2R/5$.

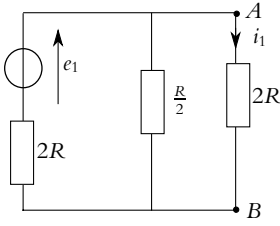


Figure S2.13

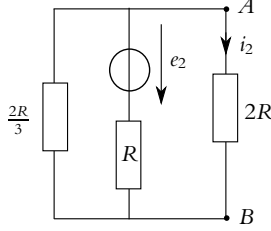


Figure S2.14

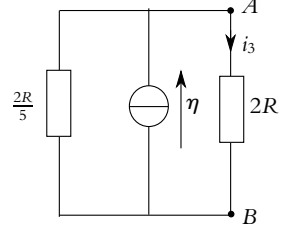


Figure S2.15

Pour le premier circuit, en associant les résistors de résistance $R/2$ et $2R$ en parallèle, on obtient le circuit de la figure (S2.16). On retrouve alors un diviseur de tension ce qui donne la tension aux bornes de ces deux résistors :

$$V_1 = V_A - V_B = e_1 \frac{\frac{2R}{5}}{\frac{2R}{5} + 2R} = \frac{e_1}{6}$$

On en déduit alors l'intensité du courant $i_1 = V_1/2R = e_1/12R$.

Pour le deuxième circuit, en associant les résistors de résistance $2R/3$ et $2R$ en parallèle, on obtient le circuit de la figure (S2.17). On retrouve alors un diviseur de tension ce qui donne la tension aux bornes de ces deux résistors :

$$V_2 = V_A - V_B = -e_2 \frac{\frac{R}{2}}{\frac{R}{2} + 2R} = -\frac{e_2}{3}$$

On en déduit alors l'intensité du courant $i_2 = \frac{V_2}{2R} = \frac{-e_2}{6R}$.

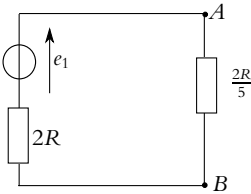


Figure S2.16

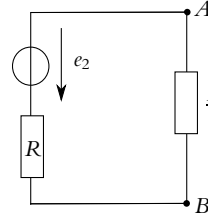


Figure S2.17

Pour le troisième circuit, un diviseur de courant donne

$$i_3 = \frac{\frac{1}{2R}}{\frac{1}{2R} + \frac{5}{2R}} \eta = \frac{1}{6} \eta$$

Finalement :

$$i = i_1 + i_2 + i_3 = \frac{e_1}{12R} - \frac{e_2}{6R} + \frac{\eta}{6}$$

ce qui correspond aux résultats obtenus par les autres méthodes.

A.8 1. Pour la partie droite du circuit la résistance R_2 est en série avec la résistance R_1 ; donc elle sont équivalentes à une résistance $R_1 + R_2$. En revanche pour la partie gauche il faut transformer le générateur de Norton en générateur de Thévenin : ce qui donne une f.e.m. $E_g = R_2 \eta$ en série avec

une résistance R_2 (figure S2.18). On peut alors associer R_1 et R_2 en série ce qui donne une résistance $R_1 + R_2$. Les deux générateurs se retrouvent alors en parallèle, il faut donc utiliser un modèle de Norton.

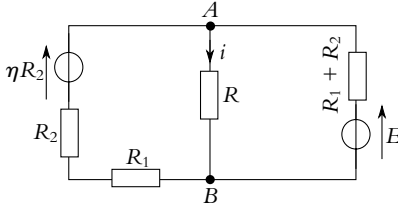


Figure S2.18

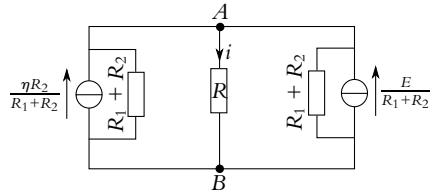


Figure S2.19

La transformation donne (figure S2.19) :

- À droite : un courant électromoteur $I_d = \frac{E}{R_1 + R_2}$ et une résistance $R_1 + R_2$ en parallèle.
- À gauche : un courant électromoteur $I_g = \frac{R_2 \eta}{R_1 + R_2}$ et une résistance $R_1 + R_2$ en parallèle.

L'association de ces deux générateurs en parallèle donne un générateur de Norton de courant électromoteur I_N et de résistance R_N :

$$I_N = I_d + I_g = \frac{E}{R_1 + R_2} + \frac{R_2 \eta}{R_1 + R_2} \quad \text{et} \quad R_N = \frac{(R_1 + R_2)^2}{2(R_1 + R_2)} = \frac{R_1 + R_2}{2}$$

Le modèle de Thévenin équivalent est une source de tension E_T et une résistance R_T en série telles que :

$$E_T = R_N I_N = \frac{E + \eta R_2}{2} \quad \text{et} \quad R_T = R_N = \frac{R_1 + R_2}{2}$$

On obtient donc les circuits équivalents de la figure (S2.20) :

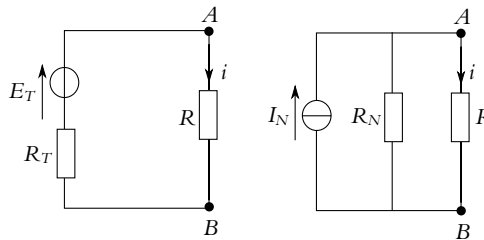


Figure S2.20

2. On peut alors au choix appliquer une loi des mailles avec le modèle de Thévenin :

$$i = \frac{E_T}{R + R_T} = \frac{E + \eta R_2}{2R + R_1 + R_2}$$

ou un diviseur de courant avec le modèle de Norton :

$$i = \frac{\frac{1}{R}}{\frac{1}{R} + \frac{1}{R_N}} I_N = \frac{E + \eta R_2}{2R + R_1 + R_2}$$

- B.1** 1. a) Les deux résistances $2R$ sont en parallèles donc équivalentes à une résistance R . Cette résistance se retrouve en série avec les deux autres, donc la résistance équivalente est $R_1 = \frac{R}{2} + \frac{R}{2} + R = 2R$.
- b) On remarque que le circuit à droite des points B_1 et B_2 est équivalent à celui de la figure (2.35). On peut donc le remplacer par un résistor de résistance $2R$. On se retrouve alors avec le même circuit que celui de la figure (2.35) et donc $R_2 = R_1 = 2R$.
- c) Ainsi ajouter un module ne change pas la valeur de la résistance équivalente entre A_1 et A_2 , donc quelle que soit le nombre de modules, on aura une résistance $R_n = 2R$.
2. a) Les deux résistances de droite sont en parallèle, donc la tension V_1 est la même que celle aux bornes de la résistance R équivalente à ces deux résistances (figure S2.21). On se retrouve alors avec un « pont diviseur » de tension et on peut écrire :

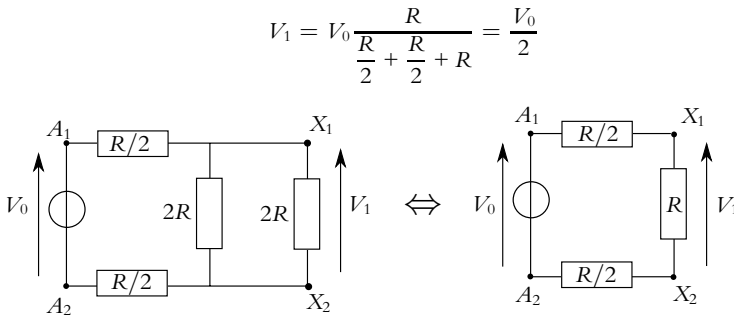


Figure S2.21

- b) On utilise une tension intermédiaire entre les points B_1 et B_2 , qui est d'ailleurs égale à la tension V_1 de la question précédente (figure S2.22). Entre B_1B_2 et X_1X_2 , comme on l'a déjà vu, on retrouve le circuit à 1 module, donc d'après le résultat précédent $V_2 = V_1/2$.

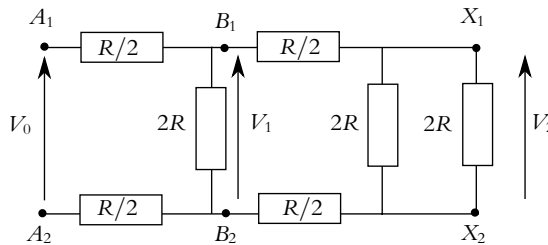


Figure S2.22

On a aussi démontré que le circuit à deux modules pouvait se ramener à un circuit à un module, et l'on se retrouve dans le cas du calcul de V_1 en fonction de V_0 . Donc finalement $V_2 = V_0/4$.

- c) Par récurrence, on déduit des résultats précédents, $V_n = \frac{V_0}{2^n}$.
- d) Lorsque n tend vers l'infini alors 2^n tend vers l'infini, donc V_∞ tend vers 0.

- B.2** 1. a) Dans le circuit en triangle, puisque $i_A = 0$, c'est le même courant qui traverse R_3 et R_2 qui se retrouvent donc en série. On obtient le circuit de gauche de la figure (S2.23). La résistance R_1

se retrouve alors en parallèle sur les deux autres. On obtient alors une résistance équivalente :

$$R_{eq} = \frac{R_1(R_2 + R_3)}{R_1 + R_2 + R_3}$$

Dans le circuit en étoile, puisque $i_A = 0$, aucun courant ne traverse R_A que l'on peut donc supprimer. On obtient alors le circuit de droite de la figure (S2.23) avec R_B et R_C en série, ce qui donne une résistance équivalente $R_B + R_C$.

Par égalité des résistances équivalentes des deux circuits, on déduit l'équation :

$$R_{eq} = \frac{R_1(R_2 + R_3)}{R_1 + R_2 + R_3} = R_B + R_C \quad (1)$$

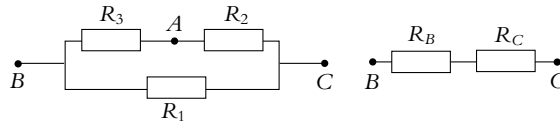


Figure S2.23

- b) Le raisonnement est le même que précédemment sauf que dans ce cas, pour le circuit triangle c'est R_1 et R_3 qui se retrouvent en série, avec R_2 en parallèle. Pour le circuit étoile, on peut supprimer R_B . Les circuits équivalents sont représentés sur la figure (S2.24).

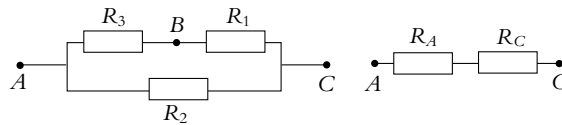


Figure S2.24

on obtient alors la relation :

$$\frac{R_2(R_1 + R_3)}{R_1 + R_2 + R_3} = R_A + R_C \quad (2)$$

- c) Encore le même raisonnement avec pour le circuit triangle, R_1 et R_2 qui se retrouvent en série, avec R_3 en parallèle. Pour le circuit étoile, on peut supprimer R_C . Les circuits équivalents sont représentés sur la figure (S2.25).

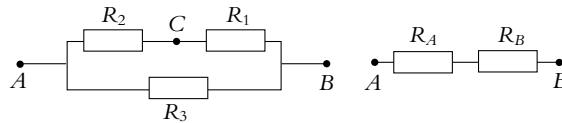


Figure S2.25

on obtient alors la relation :

$$\frac{R_3(R_1 + R_2)}{R_1 + R_2 + R_3} = R_A + R_B \quad (3)$$

On remarque qu'à partir de l'équation (1), on aurait pu procéder par permutation circulaire sur les indices (1,2,3) et (A, B, C) pour obtenir les deux autres équations.

- d) Pour obtenir R_A , on additionne les équations (2) et (3), et on soustrait l'équation (1), ce qui donne après calcul :

$$R_A = \frac{R_2 R_3}{R_1 + R_2 + R_3}$$

Les deux autres résultats sont obtenus par permutation circulaire, ou en combinant correctement les équations :

$$R_B = \frac{R_1 R_3}{R_1 + R_2 + R_3} \text{ et } R_C = \frac{R_1 R_2}{R_1 + R_2 + R_3}$$

2. a) Dans le circuit triangle, si $V_{AB} = 0$ cela revient à court-circuiter la résistance R_3 , c'est-à-dire la remplacer par un fil en reliant les points A et B . On obtient alors le circuit de gauche de la figure (S2.26) avec R_1 et R_2 en parallèle ce qui donne une conductance équivalente $G_{eq} = G_1 + G_2$. En ce qui concerne le circuit étoile, imposer $V_{AB} = 0$ revient là aussi à relier A et B par un fil : les résistances R_A et R_B se retrouvent en parallèle. On obtient le circuit de droite de la figure (S2.26). La conductance du résistor équivalent à R_A et R_B en parallèle est $G_A + G_B$ et celle du résistor équivalent à ce résistor en série avec R_C est telle que :

$$R_{eq} = R_C + \frac{1}{G_A + G_B} \Rightarrow G_{eq} = \frac{1}{R_{eq}} = \frac{G_C(G_A + G_B)}{G_A + G_B + G_C}$$

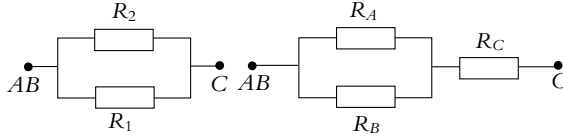


Figure S2.26

L'égalité entre les deux expressions de G_{eq} amène à l'équation suivante :

$$G_1 + G_2 = \frac{G_C(G_A + G_B)}{G_A + G_B + G_C} \quad (4)$$

- b) À l'image de ce qui a été fait précédemment, les raisonnements se ressemblent donc pour $V_{BC} = 0$ on obtient pour le circuit triangle et le circuit étoile les circuits de la figure (S2.27) ce qui amène à la relation :

$$G_2 + G_3 = \frac{G_A(G_B + G_C)}{G_A + G_B + G_C} \quad (5)$$

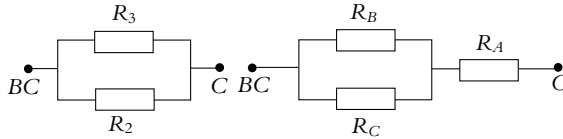


Figure S2.27

- c) Enfin pour $V_{CA} = 0$ on obtient pour le circuit triangle et le circuit étoile les circuits de la figure (S2.28) ce qui amène à la relation :

$$G_1 + G_3 = \frac{G_B(G_A + G_C)}{G_A + G_B + G_C} \quad (6)$$

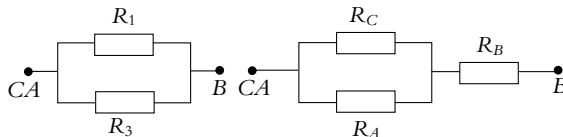


Figure S2.28

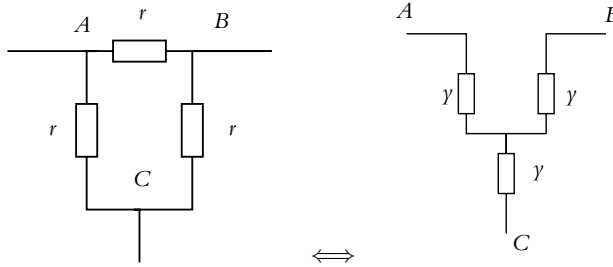
- d) Pour calculer G_1 , on combine les équations précédentes en effectuant la somme des équations (4) et (5) et la soustraction de l'équation (6). Le résultat du calcul donne :

$$G_1 = \frac{G_B G_C}{G_A + G_B + G_C}$$

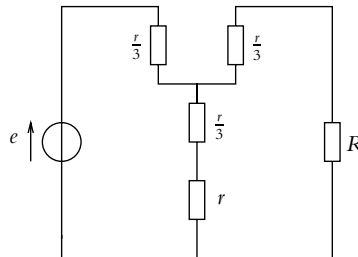
Les deux autres résultats sont obtenus par permutation circulaire, ou en combinant correctement les équations :

$$G_2 = \frac{G_A G_C}{G_A + G_B + G_C} \quad \text{et} \quad G_3 = \frac{G_A G_B}{G_A + G_B + G_C}$$

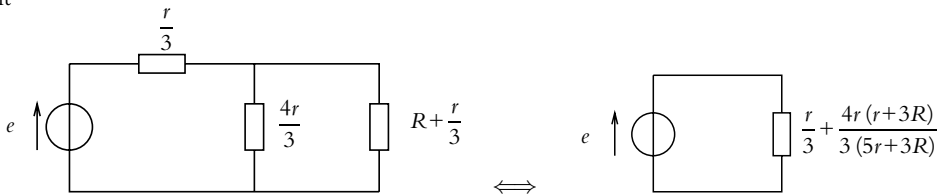
3. Par la transformation précédente, on passe du triangle ABC à l'étoile ABC :



en notant $y = \frac{rr}{r+r+r} = \frac{r}{3}$. On en déduit le circuit équivalent :



soit



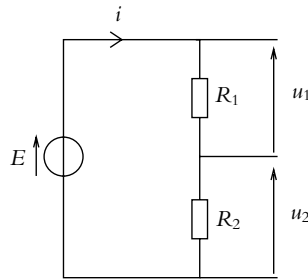
On aura l'équivalence si $\frac{r}{3} + \frac{4r(r+3R)}{3(5r+3R)} = R$ soit $r = R$.

- B.3** 1. a) Les lois de Kirchhoff sont constituées de l'ensemble des lois des nœuds $\sum_k \epsilon_k i_k = 0$ et des

lois des mailles $\sum_k \epsilon_k u_k = 0$.

- b) Une source idéale de tension est un générateur fournissant un courant d'intensité constante quelle que soit la tension à ses bornes.
c) Pour une association en série, les tensions s'ajoutent donc on effectue la « sommation » des caractéristiques selon l'axe des tensions.

- d) Le schéma d'un pont diviseur de tension est :



- e) Le courant circulant dans le circuit est $i = \frac{E}{R_1 + R_2}$ donc la tension aux bornes de la résistance s'écrit $u_1 = \frac{R_1}{R_1 + R_2} E$.
2. a) On ne peut pas appliquer la relation du pont diviseur de tension pour calculer la tension aux bornes de x car le courant circulant dans les résistances R_2 et $R_1 - x$ n'est pas nul.
- b) Le circuit comprend trois branches donc il faut *a priori* déterminer trois intensités. Ces trois intensités ne sont pas indépendantes puisqu'elles sont reliées par une loi des nœuds. Finalement il reste deux inconnues soit un système 2×2 à résoudre.
- c) L'écriture de deux lois des mailles donne

$$\begin{cases} R_1 I_1 + x I_2 = E_1 \\ x I_1 + (x + R_2) I_2 = E_2 \end{cases}$$

- d) En utilisant la méthode de substitution ou en faisant des combinaisons linéaires de ces deux équations pour éliminer I_1 puis I_2 , on obtient $I_1 = \frac{x E_2 - (x + R_2) E_1}{x^2 - R_1 (x + R_2)}$ et $I_2 = \frac{x E_1 - R_1 E_2}{x^2 - R_1 (x + R_2)}$.
- e) L'intensité dans x s'obtient alors par loi des nœuds soit

$$i = I_1 + I_2 = \frac{(x - R_1) E_2 - R_2 E_1}{x^2 - R_1 (x + R_2)}$$

- f) Si la valeur de x est telle que $I_2 = 0$, on en déduit le rapport $\frac{E_2}{E_1} = \frac{x}{R_1}$.
- g) La connaissance de x et R_1 permet donc de comparer les deux tensions E_1 et E_2 .
- h) Il est possible d'effectuer cette comparaison si les tensions sont très différentes à condition de disposer de résistances de valeurs très différentes. Les gammes de valeurs de résistance seront donc un facteur limitant pour cette comparaison.
3. a) Comme cela a été établi dans le cours, l'association en parallèle des résistances s'obtient par la relation $R_{\text{eq}} = \frac{1}{\sum_{i=1}^N \frac{1}{R_i}}$.

- b) L'absence de courant débité en dehors du dispositif permet d'écrire

$$V_1 = R_{\text{eq}} I_0 = \frac{I_0}{\sum_{i=1}^N \frac{1}{R_i}} \quad \text{soit} \quad I_0 = \left(\sum_{i=1}^N \frac{1}{R_i} \right) V_1$$

- c) L'association en série conduit à $R'_{\text{eq}} = \sum_{i=1}^N R_i$. On se reportera au cours pour la démonstration.

d) Il n'y a pas de courant débité donc $V_0 = R'_{eq} I_0 = \left(\sum_{i=1}^N R_i \right) \left(\sum_{i=1}^N \frac{1}{R_i} \right) V_1$.

e) On en déduit $\frac{V_0}{V_1} = \left(\sum_{i=1}^N R_i \right) \left(\sum_{i=1}^N \frac{1}{R_i} \right)$.

f) Si toutes les résistances prennent la valeur R , on en déduit $\frac{V_0}{V_1} = N^2$.

g) Il suffit de faire l'application numérique de la question précédente et $N = 100$.

h) On a $R - \delta R \leq R_i \leq R + \delta R$ qui permet d'écrire un encadrement des deux termes soit $NR \left(1 - \frac{\delta R}{R} \right) \leq \sum_{i=1}^N R_i \leq NR \left(1 + \frac{\delta R}{R} \right)$. L'hypothèse $\delta R \ll R$ permet d'effectuer un développement limité de

$$\frac{N}{R} \left(1 - \frac{\delta R}{R} \right) \leq \sum_{i=1}^N \frac{1}{R_i} \leq \frac{N}{R} \left(1 + \frac{\delta R}{R} \right)$$

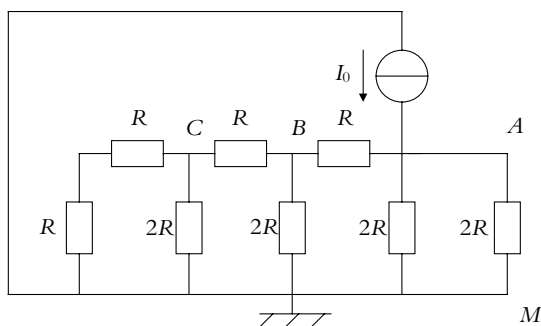
soit

$$N^2 \left(1 - 2 \frac{\delta R}{R} \right) \leq \frac{V_0}{V_1} \leq N^2 \left(1 + 2 \frac{\delta R}{R} \right).$$

i) D'après la question précédente, l'erreur maximale est $2 \frac{\delta R}{R}$.

j) L'erreur augmente moins vite que l'ordre de grandeur entre les tensions, ce qui est donc intéressant pour la comparaison.

B.4 1. On recherche la résistance x équivalente avec le schéma suivant compte tenu du fait que seul l'interrupteur K_3 est fermé.

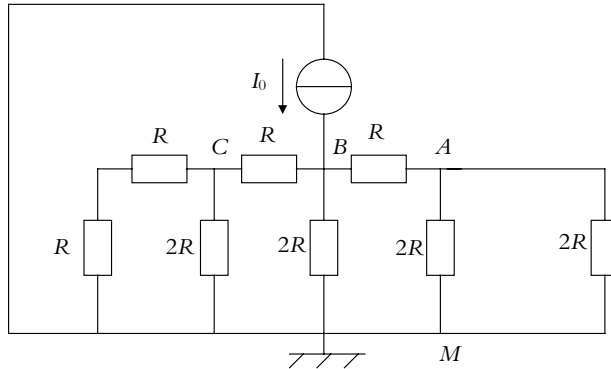


Entre M et C , l'association en série de R et R correspond à une résistance $R + R = 2R$ qui doit être associée en parallèle à une résistance $2R$ soit $R_{CM} = \frac{2R \cdot 2R}{2R + 2R} = R$.

On procède de même entre B et M et on obtient $R_{BM} = R$ puis entre A et M , ce qui permet d'en déduire que $x = R$.

2. On a donc un schéma équivalent à une association en parallèle de x et $2R$ dont la résistance équivalente est $R_{eq} = \frac{2Rx}{2R + x} = \frac{2R^2}{3R} = \frac{2}{3}R$ qu'on alimente par la source idéale de courant qui impose une intensité I_0 dans cette résistance. À ses bornes, il apparaît donc une tension $U_{AM} = R_{eq} I_0 = \frac{2}{3} R I_0$. On note que si l'interrupteur K_3 est fermé, $a_3 = 1$ et on a la tension qui vient d'être calculée et que sinon $a_3 = 0$ et la tension due à cette source idéale de courant est nulle. Par conséquent, on peut écrire que $U_{AM} = \frac{2}{3} a_3 R I_0$.

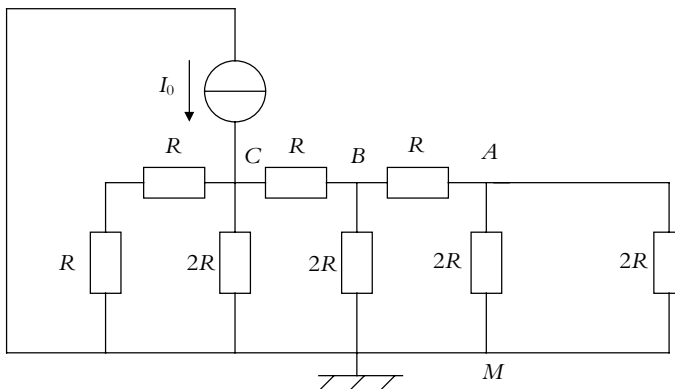
3. Le schéma est le suivant :



Entre M et C , on a la même chose que précédemment soit $R_{CM} = \frac{2R2R}{2R+2R} = R$.

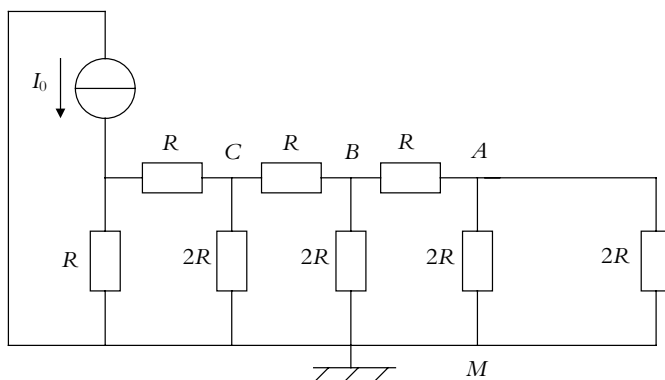
Entre M et A , l'association en parallèle de $2R$ et $2R$ donne $\frac{2R2R}{2R+2R} = R$.
Il reste à associer en parallèle $R + R = 2R$ et $R + R = 2R$ soit $x = R$.

4. Comme le schéma équivalent est le même que lorsque K_3 était seul à être fermé, on déduit la valeur $U_2 = \frac{2}{3}Ra_2I_0$ des résultats obtenus à la question 2.
5. On regarde la tension aux bornes de l'association en parallèle de $2R$ et $2R$ autrement dit d'une résistance $\frac{2R2R}{2R+2R} = R$. Il suffit alors d'appliquer la relation du pont diviseur de tension pour obtenir $U_{AM} = \frac{U_2}{2}$.
6. En utilisant les résultats des deux questions précédentes, on en déduit $U_{AM} = \frac{1}{3}Ra_2I_0$.
7. Le nouveau schéma est



Entre A et M , on a établi qu'on a R . On en déduit qu'entre B et M , on a l'association en parallèle de $2R$ et $R + R = 2R$ soit une résistance équivalente R . Par le même raisonnement, on a $x = R$ et le même montage équivalent.

8. Comme le schéma équivalent est le même que lorsque K_3 était seul à être fermé, on déduit la valeur $U_1 = \frac{2}{3} R a_1 I_0$ des résultats obtenus à la question 2.
9. On regarde la tension aux bornes de l'association en parallèle de $2R$ et $2R$ autrement dit d'une résistance $\frac{2R \cdot 2R}{2R + 2R} = R$. Il suffit alors d'appliquer la relation du pont diviseur de tension pour obtenir $U_2 = \frac{U_1}{2}$.
10. En utilisant les résultats précédents, on a $U_{AM} = \frac{U_2}{2} = \frac{U_1}{4}$.
11. En explicitant l'expression de U_1 en fonction de I_0 , on tire $U_{AM} = \frac{1}{6} R a_1 I_0$.
12. On recommence avec le schéma suivant



En effectuant le même raisonnement, on obtient le même montage équivalent avec $x = R$.

13. Comme le schéma équivalent est le même que lorsque K_3 était seul à être fermé, on déduit la valeur $U_0 = \frac{2}{3}Ra_0I_0$ des résultats obtenus à la question 2.
14. On regarde la tension aux bornes de l'association en parallèle de $2R$ et $2R$ autrement dit d'une résistance $\frac{2R \cdot 2R}{2R + 2R} = R$. Il suffit alors d'appliquer la relation du pont diviseur de tension pour obtenir $U_1 = \frac{U_0}{2}$.
15. En utilisant les résultats précédents, on a $U_{AM} = \frac{U_2}{2} = \frac{U_1}{4} = \frac{U_0}{8}$.
16. En explicitant l'expression de U_0 en fonction de I_0 , on tire $U_{AM} = \frac{1}{12}Ra_1I_0$.
17. En appliquant le principe de superposition, on obtient la tension U_{AM} comme somme des différentes expressions trouvées précédemment soit

$$U_{AM} = \frac{RI_0}{12} (a_0 + 2a_1 + 4a_2 + 8a_3)$$

18. On obtient donc l'expression demandée avec $k = 12$.
19. La valeur minimale de U_{AM} est obtenue lorsque tous les interrupteurs sont ouverts soit $U_{AM,min} = 0,00 \text{ V}$.
20. La plus petite variation possible de U_{AM} vaut $\frac{RI_0}{12} = 0,83 \text{ V}$.
21. La valeur maximale est obtenue lorsque tous les interrupteurs sont fermés soit $U_{AM,max} = 12,5 \text{ V}$.

Chapitre 3

A.1 1. En prenant les conventions de la figure (S3.1), on écrit une loi des mailles :

$$E = L \frac{di}{dt} + ri \Rightarrow \frac{di}{dt} + \frac{i}{\tau} = \frac{E}{r} \quad \text{avec} \quad \tau = \frac{L}{r}.$$

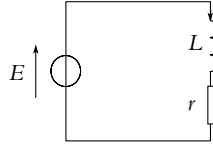


Figure S3.1

La solution donne :

$$i(t) = A \exp\left(-\frac{t}{\tau}\right) + \frac{E}{r}$$

Le premier terme avec l'exponentielle correspond au régime transitoire, et le terme E/r au régime permanent. L'intensité est continue dans la bobine et comme elle est nulle à $t = 0^-$, $i(t = 0^+) = 0$ entraîne $A = -E/r$.

2. On écrit une loi des nœuds (figure S3.2) soit $i = i_1 + i_2$ et deux lois des mailles :

$$E = ri_1 + L \frac{di_1}{dt} \quad \text{et} \quad E = u_C + Ri_2 = u_C + RC \frac{du_C}{dt}$$

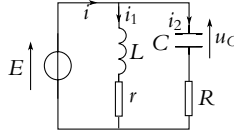


Figure S3.2

La résolution de la première équation comme dans la question précédente donne :

$$i_1(t) = A_1 \exp\left(-\frac{t}{\tau_1}\right) + \frac{E}{r} \quad \text{où} \quad \tau_1 = \frac{L}{r}$$

La résolution de la seconde équation donne :

$$u_C(t) = A_2 \exp\left(-\frac{t}{\tau_2}\right) + E \quad \text{où} \quad \tau_2 = RC$$

On en déduit :

$$i_2 = C \frac{du_C}{dt} = -\frac{CA_2}{\tau_2} \exp\left(-\frac{t}{\tau_2}\right) = -\frac{A_2}{R} \exp\left(-\frac{t}{\tau_2}\right)$$

Pour déterminer les constantes, on écrit la continuité de i_1 dans la bobine et de u_C aux bornes de C à l'instant $t = 0$, sachant que $i_1(0^-) = 0$ et $u_C(0^-) = 0$, ce qui donne : $A_1 = -E/r$ et $A_2 = -E$.

On en déduit :

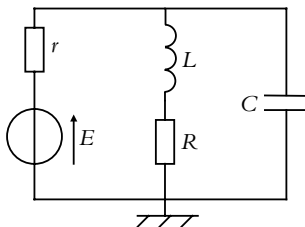
$$i = i_1 + i_2 = -\frac{E}{r} \exp\left(-\frac{t}{\tau_1}\right) + \frac{E}{r} + \frac{E}{R} \exp\left(-\frac{t}{\tau_2}\right)$$

Pour annuler le régime transitoire, il faut annuler le terme en exponentielle pour tout t :

$$-\frac{E}{r} \exp\left(-\frac{t}{\tau_1}\right) + \frac{E}{R} \exp\left(-\frac{t}{\tau_2}\right) = 0 \quad \forall t \Rightarrow r = R \quad \text{et} \quad \tau_1 = \tau_2$$

soit finalement $r = R$ et $C = L/r^2$.

A.2 Le circuit considéré est le suivant :



1. L'intensité traversant l'inductance L pour $t < 0$ s'obtient en écrivant la loi des mailles sur la maille comprenant le générateur et la bobine soit $E - r i_r - R i + L \frac{di}{dt} = 0$. Or pour $t < 0$, le circuit est en régime permanent donc $\frac{di}{dt} = 0$ et il n'y a pas de courant dans la capacité donc $i_r = i$. On en déduit $i(t < 0) = \frac{E}{r + R}$. On utilise alors la continuité de l'intensité du courant traversant une inductance L et on obtient $i(0^+) = \frac{E}{r + R}$.
2. On utilise la continuité de la tension u aux bornes de la capacité C qui est aussi la tension aux bornes de la bobine ou de l'ensemble L et R . On écrit la loi des mailles à $t = 0^+$ soit $R i(0^+) + L \frac{di}{dt}(0^+) = 0$. On en déduit $\frac{di}{dt}(0^+) = -\frac{RE}{L(r + R)}$.
3. On écrit la loi des mailles sur la maille comprenant le générateur et la bobine soit $E - r i_r - R i + L \frac{di}{dt} = 0$ ainsi que la loi des nœuds, ce qui donne $i_r = i + i_C$. Par ailleurs, la relation entre intensité et tension pour une capacité fournit $i_C = C \frac{du}{dt}$ avec $u = R i + L \frac{di}{dt}$. On en déduit $\frac{d^2 i}{dt^2} + \left(\frac{1}{rC} + \frac{R}{L} \right) \frac{di}{dt} + \frac{R + r}{rLC} i = \frac{E}{rLC}$ en éliminant i_C et u .
4. La résistance interne d'un générateur est généralement de $r = 50 \, \Omega$ (on se reportera au chapitre sur l'instrumentation à ce propos).
5. Pour obtenir un régime quasipériodique ou pseudo-périodique, il faut que le déterminant de l'équation caractéristique associée à l'équation différentielle du second ordre à coefficients constants soit négatif. Ici on a $\Delta = \left(\frac{1}{rC} + \frac{R}{L} \right)^2 - \frac{4(r + R)}{rLC}$ qui doit être négatif. On peut réécrire cette condition sous la forme

$$r^2 R^2 C^2 - 2rLC(2r + 3R) + L^2 < 0$$

Il s'agit d'un trinôme du second degré en C qui doit être négatif, il faut donc que la valeur de C soit entre les racines de ce polynôme en C pour que l'inégalité soit vérifiée. Cette condition donne $\frac{rL(2r + 3R) \pm \sqrt{r^2 L^2 (2r + 3R)^2 - r^2 R^2 L^2}}{R^2 r^2}$ soit numériquement $3,36 \, \mu\text{F} < C < 2,64 \, \text{mF}$.

6. D'après la réponse à la question précédente, le régime est pseudopériodique et la résolution en tenant compte des conditions initiales établies au début du problème, on obtient

$$i = -\frac{2rRC \exp\left(-\left(\frac{1}{2rC} + \frac{R}{2L}\right)t\right)}{(r + R) \sqrt{4(r + R)rLC - (L + rRC)^2}} \sin \frac{\sqrt{4(r + R)rLC - (L + rRC)^2}}{2rLC} t$$

- A.3** 1. Les deux condensateurs (figure S3.3) sont orientés en convention récepteur :

$$i = C_1 \frac{du_{C_1}}{dt} = C_2 \frac{du_{C_2}}{dt}$$

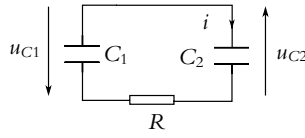


Figure S3.3

La loi des mailles donne : $u_{C1} + u_{C2} + Ri = 0$. On dérive l'équation différentielle et on utilise les relations entre i et les tensions ; on obtient alors :

$$i \left(\frac{1}{C_1} + \frac{1}{C_2} \right) + R \frac{di}{dt} = 0 \Rightarrow \frac{di}{dt} + \frac{i}{\tau} = 0$$

avec pour τ l'expression de l'énoncé. La résolution donne $i(t) = A \exp(-t/\tau)$.

Il faut maintenant déterminer la constante A . On sait que les tensions aux bornes des condensateurs sont des fonctions continues du temps, soit $u_{C1}(0^+) = Q_0/C_1$ et $u_{C2}(0^+) = 0$. La loi de mailles à $t = 0^+$ donne donc :

$$u_{C1}(0^+) + u_{C2}(0^+) + Ri(0^+) = 0 \Rightarrow i(0^+) = Q_0/R \Rightarrow A = Q_0/RC_1$$

2. Les conventions récepteur des condensateurs entraînent :

$$i = \frac{dQ_1}{dt} = \frac{dQ_2}{dt}$$

soit en intégrant :

$$Q_1 = -\frac{Q_0\tau}{RC_1} \exp\left(-\frac{t}{\tau}\right) + te_1 \quad \text{et} \quad Q_2 = -\frac{Q_0\tau}{RC_1} \exp\left(-\frac{t}{\tau}\right) + te_2$$

Or les charges sont continues dans les condensateurs donc à $t = 0^+$ $Q_1 = Q_0$ et $Q_2 = 0$. On en déduit après calcul :

$$te_1 = Q_0 \left(\frac{2C_2 + C_1}{C_1 + C_2} \right) \quad \text{et} \quad te_2 = Q_0 \left(\frac{C_1}{C_1 + C_2} \right)$$

3. La puissance reçue par un condensateur est $P = iu_C = \frac{d}{dt} \left(\frac{Q^2}{2C} \right)$, d'où l'énergie reçue pendant l'intervalle de temps dt : $\delta W = P dt$. Pour avoir l'énergie totale reçue, on intègre δW entre $t_i = 0$ et t_f :

$$W_1 = \int_0^{t_f} \frac{d}{dt} \left(\frac{Q^2}{2C_1} \right) dt = \left[\frac{Q_1^2(t)}{2C_1} \right]_0^{t_f} = \frac{1}{2C_1} [Q_1^2(t_f) - Q_1^2(0)]$$

De même, $W_2 = \frac{1}{2C_2} Q_2^2(t_f)$ puisque $Q_2(0) = 0$.

4. L'énergie reçue par la résistance pendant dt est $\delta W = Ri^2 dt$, d'où l'énergie reçue entre $t_i = 0$ et t quelconque :

$$W_R = \int_0^t Ri^2 dt = \frac{Q_0^2}{RC_1^2} \int_0^t \exp\left(-\frac{2t}{\tau}\right) dt = -\frac{Q_0^2}{RC_1^2} \frac{\tau}{2} \left[\exp\left(-\frac{2t}{\tau}\right) \right]_0^t$$

soit finalement :

$$W_R = -\frac{Q_0^2 C_2}{2C_1(C_1 + C_2)} \left(\exp\left(-\frac{2t}{\tau}\right) - 1 \right)$$

A.4 Condensateur en parallèle sur une bobine avec échelon de courant

1. On se réfère aux notations de la figure (S3.4). La loi des nœuds donne $I_0 = i_C + i$ pour $t > 0$.

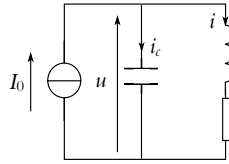


Figure S3.4

On peut aussi écrire $i_c = C \frac{du}{dt}$ et $u = L \frac{di}{dt} + Ri$. Ainsi :

$$I_0 = C \frac{du}{dt} + i = LC \frac{d^2 i}{dt^2} + RC \frac{di}{dt} + i \Rightarrow \frac{d^2 i}{dt^2} + \frac{R}{L} \frac{di}{dt} + \frac{1}{LC} i = \frac{1}{LC} I_0$$

La forme canonique est :

$$\frac{d^2 i}{dt^2} + \frac{\omega_0}{Q} \frac{di}{dt} + \omega_0^2 i = \omega_0^2 I_0$$

Par identification des termes on obtient : $\omega_0 = 1/\sqrt{LC}$ et $Q = \frac{L\omega_0}{R} = \frac{1}{R} \sqrt{\frac{L}{C}}$.

2. L'équation caractéristique est $r^2 + \frac{\omega_0}{Q}r + \omega_0^2 = 0$.

Le discriminant $\Delta = \omega_0^2 \sqrt{\frac{1}{Q^2} - 4} = \omega_0^2 \sqrt{\frac{R^2 C}{L}} - 4$ est négatif d'après la condition de l'énoncé ; on en déduit que les racines sont complexes et que la solution pour $i(t)$ est pseudo-périodique. Les racines sont donc :

$$r_1 = \frac{1}{\tau} - j\Omega \text{ et } r_2 = \frac{1}{\tau} + j\Omega \text{ avec } \frac{1}{\tau} = \frac{\omega_0}{2Q} \text{ et } \Omega = \frac{\omega_0}{2} \sqrt{4 - \frac{1}{Q^2}}$$

La solution de l'équation homogène pour $i(t)$ est $A \exp(-\frac{t}{\tau}) \cos(\Omega t + \varphi)$ à laquelle il faut ajouter la solution particulière I_0 soit finalement :

$$i(t) = A \exp(-\frac{t}{\tau}) \cos(\Omega t + \varphi) + I_0$$

3. a) Les grandeurs continues à $t = 0$ sont la tension u aux bornes du condensateur et le courant i dans la bobine.
 b) Pour $t < 0$, ces deux grandeurs sont nulles donc $i(0^+) = i(0^-) = 0$ et $u(0^+) = u(0^-) = 0$. Avec l'équation $u = L \frac{di}{dt} + Ri$ valable aussi à $t = 0^+$, on obtient $\frac{di}{dt}(0^+) = 0$.
 c) On peut alors écrire les équations suivantes :

$$\begin{cases} i(0^+) = 0 = A \cos \varphi + I_0 \\ \frac{di}{dt}(0^+) = 0 = A \left(-\Omega \sin \varphi - \frac{1}{\tau} \cos \varphi \right) \end{cases} \Rightarrow \begin{cases} A \cos \varphi = -I_0 \\ \tan \varphi = -\frac{1}{\Omega \tau} \end{cases}$$

Il faut choisir un signe pour $\cos \varphi$, par exemple positif et $\varphi \in \left[-\frac{\pi}{2}, \frac{\pi}{2}\right]$. Alors :

$$A = -\frac{I_0}{\cos \varphi} = -I_0 \sqrt{(1 + \tan^2 \varphi)} = -I_0 \sqrt{\left(1 + \frac{1}{(\Omega \tau)^2}\right)} \text{ et } \varphi = -\arctan \frac{1}{\Omega \tau}$$

Attention, si l'on avait choisi $\cos \varphi < 0$, alors on aurait $\varphi \in \left[\frac{\pi}{2}, \frac{3\pi}{2}\right]$:

$$A = -\frac{I_0}{\cos \varphi} = I_0 \sqrt{(1 + \tan^2 \varphi)} = I_0 \sqrt{\left(1 + \frac{1}{(\Omega \tau)^2}\right)} \text{ et } \varphi = -\arctan \frac{1}{\Omega \tau} + \pi$$

B.1 1. On utilise la loi d'Ohm pour une résistance donc $[R] \sim \frac{[u]}{[i]}$ et la relation tension-courant

pour une bobine $u = L \frac{di}{dt}$ d'où $[L] \sim \frac{[u][t]}{[i]}$. On en déduit que $\alpha = 1$ et $\beta = -1$.

2. a) Pour $0 \leq t \leq T/2$, la loi des mailles donne :

$$u(t) = Ri + L \frac{di}{dt} \Rightarrow \frac{di}{dt} + \frac{i}{\tau} = \frac{E}{L} \Rightarrow i(t) = A \exp\left(-\frac{t}{\tau}\right) + \frac{E}{R}$$

L'intensité est continue dans une bobine donc $i(0^+) = i(0^-)$. On suppose que pour $t < 0$ l'intensité est nulle, donc $i(0^+) = 0$ et $A = -E/R$.

b) La relation tension-courant pour une bobine donne $u_L = L \frac{di}{dt}$ d'où :

$$u_L = -L \frac{E}{R} \left(-\frac{1}{\tau}\right) \exp\left(-\frac{t}{\tau}\right) = E \exp\left(-\frac{t}{\tau}\right)$$

c) Le régime permanent correspond à $t \rightarrow \infty$, donc en régime permanent :

$$i \rightarrow \frac{E}{R} \text{ et } u_L \rightarrow 0$$

En ce qui concerne les tangentes à l'origine : $\frac{di}{dt} = \frac{u_L}{L}$ donc à $t = 0^+$:

$$\frac{di}{dt} = \frac{E}{L}$$

Pour déterminer $\frac{du_L}{dt}$, on dérive l'expression trouvée précédemment :

$$\frac{du_L}{dt} = -\frac{E}{\tau} \exp\left(-\frac{t}{\tau}\right) \Rightarrow \frac{du_L}{dt}(0^+) = -\frac{E}{\tau}$$

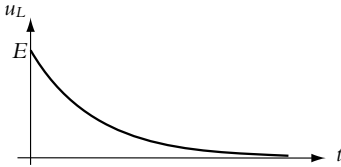


Figure S3.5

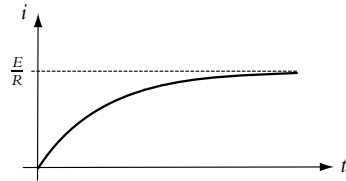


Figure S3.6

3. Pour l'intervalle de temps $\frac{T}{2} \leq t \leq T$, la loi des mailles conduit à l'équation différentielle :

$$\frac{di}{dt} + \frac{i}{\tau} = 0 \Rightarrow i(t) = A \exp\left(-\frac{t}{\tau}\right)$$

Cette fois-ci, la constante est déterminée par la continuité de i dans la bobine à l'instant $t = T/2$ soit :

$$i\left(\frac{T}{2}^+\right) = i\left(\frac{T}{2}^-\right) \Rightarrow A \exp\left(-\frac{T}{2\tau}\right) = \frac{E}{R} \left(1 - \exp\left(-\frac{T}{2\tau}\right)\right)$$

On en déduit $A = \frac{E}{R} \left(\exp\left(\frac{T}{2\tau}\right) - 1\right)$. Pour déterminer u_L , on dérive i soit :

$$u_L = L \frac{di}{dt} = -E \left(\exp\left(\frac{T}{2\tau}\right) - 1\right) \exp\left(-\frac{t}{\tau}\right)$$

4. Détermination de la période du créneau : $T = \frac{1}{f} = 1 \text{ ms}$.

Détermination de τ : $\tau = L/R = 1 \text{ ms}$.

Étant donné que $\tau = T$, les courbes exponentielles sont loin de leurs asymptotes lors des changements de valeur de $u(t)$, il est difficile de tracer qualitativement l'allure des courbes sans étude préalable. On va donc déterminer les expressions de $i(t)$ et $u_L(t)$ pour $t \leq 2T$.

Si l'on note $a = \exp(-T/2)$, sachant que :

- pour $0 \leq t \leq T/2$ et $T \leq t \leq 3T/2$ on a $u(t) = E$,
 - pour $T/2 \leq t \leq T$ et $3T/2 \leq t \leq 2T$ on a $u(t) = 0$,
- et en utilisant la continuité de $i(t)$ aux instants $T/2$, T , et $3T/2$, on obtient :

- Pour $0 \leq t \leq T/2$, $i(t) = \frac{E}{R} \left(1 - \exp\left(-\frac{t}{\tau}\right) \right)$.
- Pour $T/2 \leq t \leq T$, $i(t) = \frac{E}{R} \frac{1-a}{a} \exp\left(-\frac{t}{\tau}\right)$.
- Pour $T \leq t \leq 3T/2$, $i(t) = \frac{E}{R} \left(\frac{-a^2 + a - 1}{a^2} \exp\left(-\frac{t}{\tau}\right) + 1 \right)$.
- Pour $3T/2 \leq t \leq 2T$, $i(t) = \frac{E}{R} \frac{-a^3 + a^2 - a + 1}{a^3} \exp\left(-\frac{t}{\tau}\right)$.

Avec ces valeurs de T et τ , on obtient les courbes suivantes pour i (figure S3.7) et u_L (figure S3.8) tracées sur six périodes.

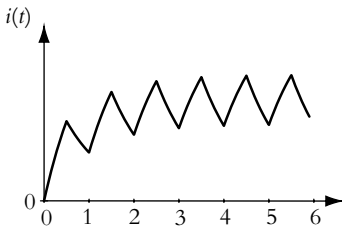


Figure S3.7

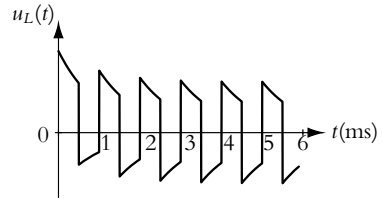
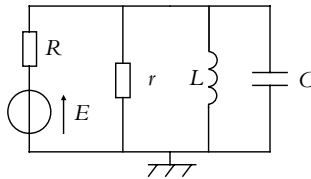


Figure S3.8

B.2 Le circuit étudié est le suivant :



1. On utilise la continuité de l'intensité du courant traversant une inductance et de la tension aux bornes d'une capacité pour déterminer les conditions initiales juste après la fermeture de l'interrupteur K .

On obtient immédiatement $i_1(0^+) = u(0^+) = 0$. Par la loi d'Ohm, on a $i_3(0^+) = \frac{u(0^+)}{r} = 0$. En écrivant la loi des nœuds avec les résultats précédents, on a $i(0^+) = i_2(0^+)$ et il suffit d'écrire une loi des mailles pour obtenir $i(0^+) = i_2(0^+) = \frac{E}{R}$.

2. En régime permanent, les grandeurs sont constantes et les dérivées temporelles nulles. Les relations intensité - tension pour les dipôles donnent $i_2(\infty) = u(\infty) = 0$. Alors en refaisant le même raisonnement qu'à la question précédente, on obtient $i_3(\infty) = 0$ et $i(\infty) = i_1(\infty) = \frac{E}{R}$.

3. On a $u = E - Ri = r i_3 = L \frac{di_1}{dt}$ et $u = \frac{du}{dt}$ soit en tenant compte du fait que E est constante, on obtient par dérivation $\frac{di}{dt} = -\frac{r}{R} \frac{di_3}{dt}$, $\frac{di_1}{dt} = \frac{r}{L} i_3$, $\frac{di_2}{dt} = rC \frac{d^2 i_3}{dt^2}$. On dérive la loi des nœuds $i = i_1 + i_2 + i_3$ et en utilisant les relations précédentes, on obtient $\frac{d^2 i_3}{dt^2} + \frac{r+R}{RrC} \frac{di_3}{dt} + \frac{1}{LC} i_3 = 0$.
4. Par identification, on en déduit $\omega_0 = \frac{1}{\sqrt{LC}}$ et $\lambda = \frac{R+r}{2Rr} \sqrt{\frac{L}{C}}$.
5. Le régime sera pseudo-périodique si le discriminant de l'équation caractéristique $r^2 + 2\lambda\omega_0 r + \omega_0^2 = 0$ est négatif soit $4\lambda^2\omega_0^2 - 4\omega_0^2 < 0$ qu'on peut écrire $\lambda < 1$ ou en explicitant la valeur de λ $\frac{R+r}{2rR} < \sqrt{\frac{C}{L}}$.
6. La grandeur λ caractérise l'amortissement des oscillations.
7. Dans le cas du régime pseudo-périodique considéré ici, la solution de l'équation caractéristique s'écrit $r_{\pm} = -\lambda\omega_0 \pm j\omega_0\sqrt{1-\lambda^2}$. On en déduit l'expression de la pseudo-pulsation $\omega = \omega_0\sqrt{1-\lambda^2} = \frac{\sqrt{4r^2R^2C - (R+r)^2L}}{2rRC\sqrt{L}}$ et de la pseudo-période $T = \frac{2\pi}{\omega} = \frac{4\pi RrC\sqrt{L}}{\sqrt{4R^2r^2C - (R+r)^2L}}$.
8. Dans ces conditions, l'intensité i_3 s'écrit $i_3 = (A \cos \omega t + B \sin \omega t) \exp(-\lambda\omega_0 t)$. La détermination des constantes A et B s'obtient à partir des conditions initiales de la question 1. soit $i_3(0) = 0$ et $\frac{di_3}{dt}(0) = \frac{E}{RrC}$.
Or $i_3(0) = A$ et $\frac{di_3}{dt} = ((-\omega A - \lambda\omega B) \sin \omega t + (\omega B - \lambda\omega A) \cos \omega t) \exp(-\lambda\omega_0 t)$ donc $\frac{di_3}{dt}(0) = \omega B - \lambda\omega A$. D'où $A = 0$, $B = \frac{E}{RrC\omega}$ et $i_3 = \frac{E}{RrC\omega} \sin \omega t \exp(-\lambda\omega_0 t)$.
9. Le décrément logarithmique est défini par $\delta = \ln \frac{i_3(t)}{i_3(t+T)}$ et son calcul conduit à $\delta = \lambda\omega_0 T$.
10. Le régime est établi avec une précision de un millièmètre si l'amplitude de i_3 est inférieure au millièmètre de sa valeur maximale soit $\frac{E}{1000RrC\omega}$. On résout donc $i_3 = \frac{E}{RrC\omega} \exp(-\lambda\omega_0 t) < \frac{E}{1000RrC\omega}$. On obtient $t \geq \frac{\ln 1000}{\lambda\omega_0}$.
11. Pour déterminer les nouvelles conditions initiales juste après l'ouverture de l'interrupteur, on raisonne comme à la question 1. On obtient $i = u = i_3 = 0$ et $i_1 = \frac{E}{R} = -i_2$.
12. De même, avec un raisonnement analogue à celui de la question 3., on a $\frac{d^2 i_3}{dt^2} + \frac{1}{rC} \frac{di_3}{dt} + \frac{1}{LC} i_3 = 0$.
13. Par identification, on obtient $\omega_0 = \frac{1}{\sqrt{LC}}$ et $\lambda = \frac{1}{2r} \sqrt{\frac{L}{C}}$.
14. Comme $\frac{1}{2r} < \frac{r+R}{RrC} < \sqrt{\frac{C}{L}}$, on a un régime pseudo-périodique.
15. On aura un signal non nul uniquement si le courant dans la bobine est non nul.

B.3 1. On note i l'intensité traversant C_b et R_b orientée dans le sens M vers N . On peut alors écrire $V(t) = V_M - V_N + V_N - V_P$, avec $V_M - V_N = R_b i$ et $V_N - V_P = u(t)$. D'autre part la relation tension-courant pour C_b donne $i = -C_b \frac{dV}{dt}$ puisque C_b est orientée en convention générateur. Ainsi :

$$V = R_b \left(-C_b \frac{dV}{dt} \right) + u \Rightarrow \frac{u - V}{R_b} = C_b \frac{dV}{dt}$$

On écrit maintenant une loi des nœuds en $N : i = i_1 + i_2$. Or pour R_a la loi d'Ohm donne $u = R_a i_2$ et pour C_a , on obtient $i_1 = C_a \frac{du}{dt}$, d'où :

$$i = i_1 + i_2 = C_a \frac{du}{dt} + \frac{u}{R_a}$$

2. On dérive la première des équations précédentes et on utilise la relation $i = -C_b \frac{dV}{dt}$:

$$C_b \frac{d^2 V}{dt^2} = \frac{1}{R_b} \frac{du}{dt} - \frac{1}{R_b} \frac{dV}{dt} \Rightarrow -\frac{di}{dt} = \frac{1}{R_b} \frac{du}{dt} + \frac{1}{R_b C_b} i$$

On utilise la deuxième équation ainsi que son expression dérivée pour éliminer i et $\frac{di}{dt}$ dans l'expression précédente. Il vient :

$$-C_a \frac{d^2 u}{dt^2} - \frac{1}{R_a} \frac{du}{dt} = \frac{1}{R_b} \frac{du}{dt} + \frac{1}{R_b C_b} \left(C_a \frac{du}{dt} + \frac{u}{R_a} \right)$$

soit en regroupant les termes :

$$\frac{d^2 u}{dt^2} + \left(\frac{1}{R_a C_a} + \frac{1}{R_b C_b} \right) \frac{du}{dt} + \frac{1}{R_a C_a R_b C_b} u = 0$$

On obtient bien l'équation demandée avec :

$$\omega_0^2 = \frac{1}{R_a C_a R_b C_b} \text{ et } \frac{\omega_0}{Q} = \frac{1}{R_a C_a} + \frac{1}{R_b C_b} \Rightarrow \omega_0 = \frac{1}{\sqrt{R_a C_a R_b C_b}}$$

et

$$Q = \frac{\sqrt{R_a C_a R_b C_b}}{R_a C_a + R_b C_b}$$

3. a) Le discriminant de l'équation caractéristique est $\Delta = \omega_0^2 \left(\frac{1}{Q^2} - 4 \right)$, or :

$$\begin{aligned} \frac{1}{Q^2} - 1 &= \frac{(R_a C_a + R_b C_b)^2}{R_a C_a R_b C_b} - 4 = \frac{(R_a C_a + R_b C_b)^2 - 4 R_a C_a R_b C_b}{R_a C_a R_b C_b} \\ &= \frac{(R_a C_a - R_b C_b)^2}{R_a C_a R_b C_b} \end{aligned}$$

Si $R_a C_a \neq R_b C_b$ alors $\Delta > 0$ et l'équation caractéristique admet deux solutions réelles r_1 et r_2 . Le produit des racines est $r_1 r_2 = \omega_0^2 > 0$ donc les racines sont de même signe. La somme des racines est $r_1 + r_2 = -\frac{\omega_0}{Q} < 0$ donc les racines sont toutes les deux négatives.

- b) L'expression générale de $u(t)$ est donc :

$$u(t) = A_1 e^{r_1 t} + A_2 e^{r_2 t}$$

où A_1 et A_2 sont deux constantes dépendant des conditions initiales.

- c) À $t = 0$, les tensions aux bornes des condensateurs sont continues, donc $V(0^+) = V(0^-) = V_0$ et $u(0^+) = u(0^-) = 0$. Or on peut écrire $i(0^+) = u(0^+)/R_a$, donc $i(0^+) = 0$. Puisque $V = R_b i + u$, à $t = 0^+$ on a $i = V_0/R_b$ et comme $i_2 = 0$, alors $i_1 = i = V_0/R_b$.

- d) On écrit les conditions initiales pour u et $\frac{du}{dt} = i_1/C_a$:

$$\begin{cases} u(0^+) = 0 = A_1 + A_2 \\ \frac{du}{dt}(0^+) = A_1 r_1 + A_2 r_2 = \frac{V_0}{R_b C_a} \end{cases} \Rightarrow \begin{cases} A_1 = \frac{1}{r_1 - r_2} \frac{V_0}{R_b C_a} \\ A_2 = \frac{1}{r_2 - r_1} \frac{V_0}{R_b C_a} \end{cases}$$

On peut d'autre part calculer r_1 et r_2 :

$$r_1 = -\frac{\omega_0}{2Q} + \frac{\sqrt{\Delta}}{2} \quad r_2 = -\frac{\omega_0}{2Q} - \frac{\sqrt{\Delta}}{2} \Rightarrow r_1 - r_2 = \sqrt{\Delta} = \omega_0 \frac{|R_a C_a - R_b C_b|}{\sqrt{R_a C_a R_b C_b}}$$

d'où

$$A_1 = \frac{V_0}{\omega_0 |R_a C_a - R_b C_b|} \sqrt{\frac{R_a C_b}{R_b C_a}} \text{ et } A_2 = -A_1$$

- B.4** 1. a) Au bout d'un temps très long, le condensateur est chargé et l'intensité est nulle. La tension u_C (figure S3.9) aux bornes de C est égale à E et $u_L = 0$.

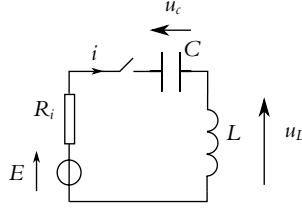


Figure S3.9

- b) Les deux grandeurs continues sont :

- la tension aux bornes de C , soit $u_C(0^+) = u_C(0^-) = 0$;
- l'intensité traversant L , soit $i(0^+) = i(0^-) = 0$.

Une loi des mailles à $t = 0^+$ donne :

$$E = R_i i + u_C + u_L \Rightarrow u_L(0^+) = E$$

- c) La relation tension-courant pour une bobine est : $u_L = L \frac{di}{dt}$ et pour un condensateur : $i = C \frac{du_C}{dt}$
On écrit de nouveau une loi des mailles que l'on dérive deux fois :

$$E = R_i i + u_C + u_L \Rightarrow 0 = R_i \frac{di}{dt} + \frac{du_C}{dt} + \frac{du_L}{dt}$$

Soit en remplaçant i et u_C :

$$0 = \frac{R_i}{L} \frac{du_L}{dt} + \frac{1}{LC} u_L + \frac{du_L}{dt}$$

- d) En considérant $R_i = 0$ et en posant $\omega_0 = 1/\sqrt{LC}$, l'équation devient une équation d'oscillateur harmonique $\ddot{u}_L + \omega_0^2 u_L = 0$, dont la solution est :

$$u_L = A \cos \omega_0 t + B \sin \omega_0 t$$

Sachant qu'à $t = 0^+$, $u_L = E$ on déduit $A = E$. L'autre condition initiale pour $\frac{du_L}{dt}(0^+)$ est plus difficile à obtenir. On repart de la loi des mailles (avec $R_i = 0$) que l'on dérive, et on calcule l'expression à $t = 0^+$.

$$E = u_C + u_L \Rightarrow 0 = \frac{du_C}{dt} + \frac{du_L}{dt} \Rightarrow 0 = \frac{i(0^+)}{C} + \frac{du_L}{dt}(0^+) \Rightarrow \frac{du_L}{dt}(0^+) = 0$$

On en déduit $B = 0$.

- e) On ne néglige plus R_i . L'équation caractéristique et son discriminant sont :

$$r^2 + r \frac{R_i}{L} + \frac{1}{LC} = 0 \text{ et } \Delta = \left(\frac{R_i}{L} \right)^2 - \frac{4}{LC} = -5.10^{11} \text{ s}^{-2}$$

Le régime est donc pseudo-périodique de temps caractéristique $\tau = L/R_i$. La décroissance de u_L à partir de la valeur initiale E va être due à une exponentielle : $\exp(-t/\tau)$. On choisit par exemple de calculer le temps au bout duquel $u_L = 1\% u_L(0^+)$, ce qui revient à chercher t tel que $\exp(-t/\tau) = 0,01$ soit $t = -\tau \ln 0,01 = 0,33 \tau$.

2. a) Lorsque le condensateur est chargé, le courant ne circule plus donc $u = U_0$. L'énergie dans le condensateur est alors $W = \frac{1}{2} C' U_0^2$ soit $W = 6,75 \text{ J}$.
- b) On transforme le générateur de Thévenin (U_0, R'_i) en modèle de Norton ($\frac{U_0}{R'_i}, R'_i$). Les deux résistances se retrouvent alors en parallèle ce qui donne une résistance équivalente $R_e = \frac{R'_i R_f}{R'_i + R_f}$. On effectue de nouveau une transformation Norton-Thévenin pour obtenir le circuit de la figure (S3.10) avec $E_t = R_e \frac{U_0}{R'_i}$ soit $E_t = \frac{R_f U_0}{R_f + R'_i}$ et $R_t = R_e$.

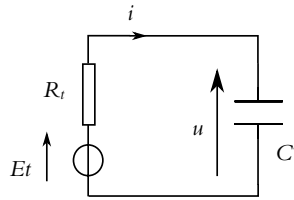


Figure S3.10

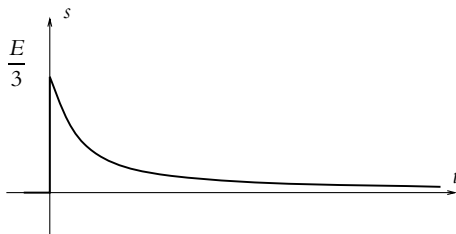
Application numérique : $E_t = 15,8 \text{ V}$ et $R_t = 9,5 \Omega$.

- c) La loi des mailles donne alors $E_t = R_t i + u = R_t C' \frac{du}{dt} + u$. La solution de l'équation est $u(t) = A \exp(-t/\tau) + E_t$ avec $\tau = R_t C'$, où A est déterminé par les conditions initiales. À $t = 0$, $u = U_0$ d'où finalement $u(t) = (U_0 - E_t) \exp(-t/\tau) + E_t$.
- d) L'ordre de grandeur de décharge du condensateur est $\tau = R_t C' = 1,4 \text{ ms}$.

B.5 1. Circuit simple :

- a) L'écriture de la loi des mailles donne $e = Ri + s$ et celle de la loi des nœuds $i = i_R + i_L$. On la continuité de l'intensité traversant une inductance permet d'écrire la continuité de i_L soit $i_L(0^+) = 0$. Cette valeur implique dans la loi des nœuds que $i(0^+) = i_R(0^+)$. On obtient alors la tension s par la relation du pont diviseur de tension $s(0^+) = \frac{R}{2} i_R(0^+) = \frac{R}{2} i(0^+)$. Finalement en reportant dans la loi des mailles, on en déduit $E = Ri(0^+) + s(0^+)$ soit $i(0^+) = \frac{2E}{3R}$. Par conséquent, l'intensité i est discontinue puisque $i(0^-) = 0$.
- b) D'après la question précédente, on a $s(0^+) = \frac{E}{3}$ qui est également discontinue puisque $s(0^-) = 0$.
- c) Lorsque le temps t tend vers l'infini, on obtient le régime permanent et les dérivées temporelles sont donc nulles. On obtient donc $s = L \frac{di}{dt} = 0$ pour un temps infini.
- d) La loi des mailles $e = Ri + s$ avec $s = L \frac{di_L}{dt}$ et la loi des nœuds $i = i_L + i_R$ permet d'obtenir en dérivant $\frac{de}{dt} = R \frac{di_L}{dt} + R \frac{di_R}{dt} + \frac{ds}{dt} = R \frac{s}{L} + R \frac{d}{dt} \left(\frac{2s}{R} \right) + \frac{ds}{dt}$. L'équation différentielle demandée est donc $3 \frac{ds}{dt} + \frac{R}{L} s = \frac{de}{dt}$.
- e) On résout cette équation différentielle dont la solution est la somme de la solution de l'équation homogène associée $S \exp\left(-\frac{R}{3L}t\right)$ et d'une solution particulière $s = 0$. On détermine la constante d'intégration S en utilisant la condition initiale soit finalement $s = \frac{E}{3} \exp\left(-\frac{R}{3L}t\right)$.

- f) L'allure de $s(t)$ est donc la suivante :



- g) Le temps t_0 cherché s'obtient en résolvant $s(t_0) = \frac{s(0^+)}{10}$ dont la solution est $t_0 = \frac{3L \ln 10}{R}$.
- h) On branche l'oscilloscope entre la masse et le point reliant R , L et $\frac{R}{2}$. On utilise le décalibrage vertical de l'oscilloscope de manière à placer le maximum de s soit S sur la ligne 100 % et la valeur minimale de s soit 0, 0 sur la ligne 0 % de l'écran de l'oscilloscope. On recherche ensuite la valeur de t_0 en déplaçant le curseur vertical jusqu'à ce qu'il soit sur l'intersection de la tension avec la ligne 10 %.
- i) D'après ce qui précède, on déduit $L = \frac{Rt_0}{3 \ln 10} = 4,3 \cdot 10^{-4}$ H.
- j) Pour mesurer t_0 , il faut laisser le temps au régime transitoire de se terminer, ce qui impose $\frac{T}{2} \gg t_0$ et en termes de fréquences $f \ll \frac{1}{2t_0} \simeq 10^5$ Hz.
2. Circuit plus complexe :
- a) On procède comme précédemment en appliquant la continuité de l'intensité du courant traversant les inductances. Cela impose de façon immédiate $i_4(0^+) = 0$. Par ailleurs, à $t = 0^-$, le circuit est en régime permanent donc les dérivées temporelles sont nulles, ce qui donne pour la tension aux bornes de l'inductance parcourue par le courant i_1 $u(0^-) = L \frac{di_1}{dt}(0^-) = 0$. On en déduit par la loi d'Ohm que $i_2(0^-) = \frac{u(0^-)}{R} = 0$ puis que $i(0^-) = i_1(0^-) = i_1(0^+)$ par continuité de l'intensité du courant traversant l'inductance. L'écriture de la loi des mailles donne $E = Ri(0^-) + Ri_3(0^-) = 2Ri(0^-)$ en utilisant la loi des nœuds $i(0^-) = i_3(0^-) + i_4(0^-) = i_3(0^-)$ soit $i(0^-) = \frac{E}{2R} = i_1(0^-)$. La continuité de l'intensité dans l'inductance permet alors d'écrire $i_1(0^+) = \frac{E}{2R}$. La loi des nœuds $i(0^+) = i_1(0^+) + i_2(0^+) = \frac{E}{2R} + i_2(0^+) = i_3(0^+) + i_4(0^+) = i_3(0^+)$ donne en reportant dans la loi des mailles $E = Ri(0^+) + Ri_2(0^+) + Ri_3(0^+)$ $i(0^+) = \frac{E}{2R}$ soit avec relations précédentes $i_2(0^+) = 0$ et $i_3(0^+) = \frac{E}{2R}$.
- b) Pour un temps infini, le régime est permanent donc les dérivées temporelles sont nulles. On en déduit $s(\infty) = 0$ et $u(\infty) = Ri_2(\infty) = 0$. En reportant dans la loi des mailles, on en déduit $E = Ri(\infty) + Ri_2(\infty) + Ri_3(\infty)$ soit $i(\infty) = \frac{E}{R}$. Alors avec la loi des nœuds $i(\infty) = i_1(\infty) + i_2(\infty)$, on a par la loi des mailles $i_1(\infty) = \frac{E}{R}$ et $i_3(\infty) = \frac{s(\infty)}{R} = 0$. Enfin la loi des nœuds donne $i(\infty) = i_3(\infty) + i_4(\infty)$ et $i_4(\infty) = \frac{E}{R}$.
- c) Par la loi d'Ohm, on a $s = Ri_3$ et la relation entre l'intensité et la tension pour une inductance donne $s = L \frac{di_4}{dt}$.
- d) La loi des nœuds s'écrit $i = i_3 + i_4$ soit en dérivant et en utilisant les relations de la question précédente, on en déduit $\frac{di}{dt} = \frac{1}{R} \frac{ds}{dt} + \frac{s}{L}$

- e) L'égalité des tensions pour deux dipôles en parallèle s'écrit $Ri_2 = L \frac{di_1}{dt}$ avec $i = i_1 + i_2$ par la loi des nœuds. On a donc $\frac{di}{dt} = \frac{R}{L}i_2 + \frac{di_2}{dt}$.
- f) Par une loi des mailles, on a $E = Ri + Ri_2 + Ri_3$, ce qui permet en dérivant et en utilisant les relations précédentes en dérivant et en utilisant les relations établies auparavant d'obtenir $\frac{dE}{dt} = 2\frac{ds}{dt} + \frac{R}{L}s + R\frac{di_2}{dt} = 3\frac{ds}{dt} + 2\frac{R}{L}s - \frac{R^2}{L}i_2$.
- g) En dérivant à nouveau cette relation et en utilisant encore les relations établies aux questions précédentes, on en déduit $\frac{d^2E}{dt^2} + \frac{R}{L}\frac{dE}{dt} = 3\frac{d^2s}{dt^2} + 4\frac{R}{L}\frac{ds}{dt} + \left(\frac{R}{L}\right)^2 s$.
- h) Ici la tension E est constante donc sa dérivée est nulle et on a l'équation différentielle

$$3\frac{d^2s}{dt^2} + 4\frac{R}{L}\frac{ds}{dt} + \left(\frac{R}{L}\right)^2 s = 0$$

- i) L'équation caractéristique de cette équation différentielle du second ordre est $3r^2 + 4\frac{R}{L}r + \left(\frac{R}{L}\right)^2 = 0$. Son discriminant s'écrit $\Delta = \left(2\frac{R}{L}\right)^2 > 0$. La solution est donc de la forme $s(t) = A \exp\left(-\frac{R}{L}t\right) + B \exp\left(-\frac{R}{3L}t\right)$.
- j) Par les conditions initiales, on a $s(0^+) = Ri_3(0^+) = \frac{E}{2}$. Or par la forme de la solution, on a $s(0^+) = A + B$. La relation demandée est donc $A + B = 0$.

B.6 1. a) En position (1) le circuit est fermé sur le générateur de f.e.m. E . La loi des mailles entraîne :

$$E = L \frac{di}{dt} + V_c$$

- b) L'énoncé permet de supposer que la tension V_c est constante. L'intégration de l'équation précédente donne alors :

$$i(t) = \frac{E - V_c}{L}t + cte$$

Sachant qu'à $t = 0$, l'intensité est égale à I_m , on obtient :

$$i(t) = \frac{E - V_c}{L}t + I_m$$

- c) En appliquant l'expression précédente au temps $t = \alpha T$, on obtient :

$$i(\alpha T) = I_m = \frac{E - V_c}{L}\alpha T + I_m$$

2. a) Le générateur est maintenant remplacé par un fil, d'où l'équation :

$$0 = L \frac{di}{dt} + V_c$$

- b) En considérant V_c constant et comme i est continue dans la bobine à l'instant αT , on obtient :

$$i(t) = -\frac{V_c}{L}t + cte \text{ avec } i(\alpha T) = I_m \Rightarrow i(t) = I_m - \frac{V_c}{L}(t - \alpha T)$$

et en remplaçant I_m par son expression :

$$i(t) = \frac{E\alpha T}{L} + I_m - \frac{V_c}{L}t$$

- c) Le régime périodique impose $I_m = i(0) = i(T)$ soit :

$$\frac{E\alpha T}{L} + I_m - \frac{V_c}{L}T = I_m \Rightarrow V_c = \alpha E$$

3. a) La relation tension-courant pour le condensateur donne $i_c = C \frac{dV_c}{dt}$ si i_c est l'intensité dans C avec une convention récepteur. Si l'on considère V_c constante alors $i_c = 0$.
Si cependant on tient compte des (légères) variations de V_c , sachant que c'est une fonction périodique de période T , on peut alors écrire :

$$\langle i_c \rangle = \frac{C}{T} \left[\int_0^{\alpha T} \frac{dV_c}{dt} dt + \int_{\alpha T}^T \frac{dV_c}{dt} dt \right]$$

soit après intégration :

$$\langle i_c \rangle = \frac{C}{T} (V_c(\alpha T) - V_c(0) + V_c(T) - V_c(\alpha T)) = V_c(T) - V_c(0) = 0$$

On note i_R l'intensité dans R. La loi des nœuds entraîne $i = i_C + i_R$ donc $\langle i \rangle = \langle i_C \rangle + \langle i_R \rangle$, soit $\langle i \rangle = \langle i_R \rangle = \langle V_c \rangle / R$, or V_c est considérée constante donc $\langle V_c \rangle = V_c$; ainsi $\langle i \rangle = V_c / R$.

On peut d'autre part, calculer directement $\langle i(t) \rangle$. On sait que c'est une fonction triangle passant de I_m à I_m entre $t = 0$ et $t = \alpha T$ et de I_m à I_m entre $t = \alpha T$ et T . Sur l'intervalle $[0, \alpha T]$, on peut écrire $i(t)$ sous la forme $i(t) = \frac{I_m - I_m}{\alpha T} t + I_m$. Le calcul de l'intégrale donne :

$$\frac{1}{T} \int_0^{\alpha T} \left(\frac{I_m - I_m}{\alpha T} t + I_m \right) dt = \frac{1}{T} \left(\frac{I_m - I_m}{\alpha T} \frac{(\alpha T)^2}{2} + I_m \alpha T \right) = \alpha \frac{I_m + I_m}{2}$$

De même sur $[\alpha T, T]$, on trouverait $(1 - \alpha) \frac{I_m + I_m}{2}$. en sommant les deux résultats, on obtient :

$$\langle i(t) \rangle = \frac{I_m + I_m}{2}$$

Des deux résultats pour $\langle i \rangle$, on déduit :

$$\frac{I_m + I_m}{2} = \frac{V_c}{R} = \frac{\alpha E}{R}$$

- b) On a donc deux équations pour I_m et I_m :

$$\begin{cases} I_m + I_m = \frac{2\alpha E}{R} \\ I_m - I_m = \frac{E(1 - \alpha)}{L} \alpha T \end{cases} \Rightarrow \begin{cases} I_m = \frac{\alpha E}{R} + \frac{E(1 - \alpha)}{2L} \alpha T \\ I_m = \frac{\alpha E}{R} - \frac{E(1 - \alpha)}{2L} \alpha T \end{cases}$$

4. Pour que $I_m > 0$, l'expression déterminée précédemment impose $R < \frac{2L}{(1 - \alpha)T}$.

Partie II – Mécanique I

Chapitre 4

- A.1** 1. Le mouvement est rectiligne à accélération constante $\vec{a}_0$. On choisit l'axe Ox (vecteur $\vec{u}_x$) comme axe du mouvement et le point O comme point de départ. L'intégration de l'accélération donne :

$$\vec{a} = \vec{a}_0 = a_0 \vec{u}_x \Rightarrow \vec{v} = (a_0 t + cte) \vec{u}_x$$

or la vitesse initiale ($t = 0$) est nulle donc $cte = 0$. L'abscisse x est obtenue par :

$$x = \frac{a_0 t^2}{2} + cte' \quad \text{avec } x(0) = 0 \quad \text{donc } cte' = 0$$

Si, on note $t_D = 26,6$ s, $a_0 = 2D/t_D^2 = 2 \times 180/26,6^2 = 0,509 \text{ m.s}^{-2}$. On obtient ensuite la vitesse à $t = t_D$: $v_D = v_0 t_D = 13,5 \text{ m.s}^{-1}$.

2. Pour la phase de freinage, on peut changer d'origine ; la nouvelle vitesse initiale est alors v_D . On note $a_f = -7 \text{ m.s}^{-2}$ l'accélération lors de cette phase ($a_f < 0$ car il y a freinage). Les intégrations de l'accélération amènent à :

$$v = \dot{x} = a_f t + v_D \text{ et } x = \frac{a_f t^2}{2} + v_D t$$

L'arrêt a lieu pour $v = 0$ soit $t = v_D/a_f = 1,93 \text{ s}$ d'où une distance de 13 m.

- A.2** 1. Soit Ox l'axe le long duquel a lieu le mouvement.

La vitesse de la voiture est $\dot{x} = v_0$ et sa position $x = v_0 t$ en prenant l'origine à la position de la voiture ou celle de la moto (ce sont les mêmes) à $t = 0$.

Quant à la moto, son accélération est $\ddot{x} = a_0$ dont on déduit par intégration sa vitesse $\dot{x} = a_0 t$ et sa position $x = \frac{1}{2} a_0 t^2$.

On cherche l'instant pour lequel les positions de la voiture et de la moto sont à nouveau identiques, ce qui revient à résoudre $v_0 t = \frac{1}{2} a_0 t^2$. Les solutions sont $t = 0$ qui n'est bien évidemment pas la solution

cherchée et $t_0 = \frac{2v_0}{a_0} = 22,2 \text{ s}$. On note qu'il faut pour faire l'application numérique transformer la vitesse en m.s^{-1} soit $v_0 = 100 \text{ km.h}^{-1} = 27,8 \text{ m.s}^{-1}$ et déterminer l'accélération par le fait qu'au bout de 10 s, la moto atteint une vitesse de 90 km.h^{-1} ou 25 m.s^{-1} soit $a_0 = \frac{v_0}{t} = 2,5 \text{ m.s}^{-2}$.

2. La distance parcourue est alors $d = v_0 t_0 = 617 \text{ m}$. On note qu'on obtient le même résultat avec $d = \frac{1}{2} a_0 t_0^2$.

3. La moto possède alors une vitesse $v_m = a_0 t_0 = 55 \text{ m.s}^{-1} = 198 \text{ km.h}^{-1}$.

- A.3** 1. La Terre effectue un tour sur elle-même en une période moyenne $T = 24 \text{ h}$ soit $T = 86\,400 \text{ s}$. Sa vitesse angulaire est $\Omega = 2\pi/T = 7,27 \cdot 10^{-5} \text{ rad.s}^{-1}$.

2. On note $\vec{r} = \overrightarrow{TM} = r\vec{u}_r$ le vecteur donnant la position du satellite en coordonnées polaires (r, θ). La période du satellite est celle de rotation de la Terre donc sa vitesse angulaire $\dot{\theta} = 2\pi/T$ est égale à $\Omega = 2\pi/T$. Il faut donc déterminer $\dot{\theta}$ connaissant l'accélération. L'expression générale de l'accélération est :

$$\vec{a} = (\ddot{r} - r\dot{\theta}^2)\vec{u}_r + (2\dot{r}\dot{\theta} + r\ddot{\theta})\vec{u}_\theta$$

Or le satellite M effectue un mouvement circulaire uniforme de centre, le centre T de la Terre, soit $\dot{r} = 0$, $\ddot{r} = 0$ et $\dot{\theta} = \Omega$, donc $a = -r\Omega^2$. Avec l'expression de l'énoncé, on déduit :

$$\Omega = \dot{\theta} = \sqrt{\frac{g_0 R^2}{r^3}} \Rightarrow r = \left(\frac{g_0 R^2}{\Omega^2} \right)^{1/3} = 42\,200 \text{ km}$$

On retient en général 36 000 km pour l'orbite géostationnaire ; il s'agit de l'altitude à partir du sol, obtenue en soustrayant le rayon de la Terre au résultat précédent.

- A.4** 1. a) On écrit x et y à $t = 0$ soit :

$$x(0) = \alpha \cos \varphi \text{ et } y(0) = \beta \cos \psi$$

or l'énoncé précise que $x(0) = a$ et $y(0) = 0$; on en déduit $\alpha = a$, $\varphi = 0$ et $\psi = 0$.

- b) Si on remplace x et y par leurs expressions dans l'équation de l'ellipse, on obtient :

$$\cos^2 \omega t + \left(\frac{\beta}{b} \right)^2 \sin^2 \omega t = 1$$

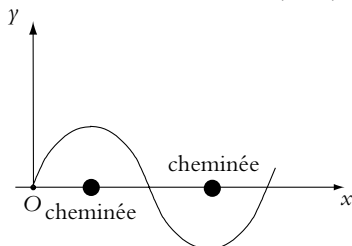
La seule possibilité pour que l'équation soit vérifiée quel que soit t est $\beta = b$.

2. a) Pour obtenir la vitesse et l'accélération on dérive par rapport au temps, sachant que $\omega = cte$:

$$\begin{cases} \dot{x} = -a\omega \sin \omega t \\ \dot{y} = a\omega \cos \omega t \end{cases} \quad \text{et} \quad \begin{cases} \ddot{x} = -a\omega^2 \cos \omega t \\ \ddot{y} = -a\omega^2 \sin \omega t \end{cases}$$

b) On remarque que $\ddot{x} = -a\omega^2 x$ et $\ddot{y} = -a\omega^2 y$, donc $\vec{a} = -k\vec{OM}$ avec $k = a\omega^2$.

A.5 On choisit un repère cartésien (O, x, y) et on note $\vec{u}_x$ et $\vec{u}_y$ les vecteurs de base. L'équation de la trajectoire, telle qu'elle apparaît sur la figure, s'écrit $y = a \sin\left(\frac{2\pi x}{L}\right)$. Il faut déterminer a .



1. Selon Ox , $\dot{x} = v_0$, soit $x = v_0 t$ puisqu'à $t = 0$, $x = 0$. Le véhicule revient sur l'axe après la sixième cheminée, ayant parcouru une distance $D = 6L = 1200$ m, la vitesse est donc $v_0 = 6L/t_f = 100 \text{ m.s}^{-1}$ ou 360 km.h^{-1} .
2. En remplaçant x par son expression dans y , on obtient $y = a \sin\left(\frac{2\pi v_0 t}{L}\right)$. L'accélération a pour composantes $\ddot{x} = \dot{v}_0 = 0$ et $\ddot{y}$. Il vient :

$$\dot{y} = a \frac{2\pi v_0}{L} \cos \frac{2\pi v_0 t}{L} \quad \text{et} \quad \ddot{y} = -a \left(\frac{2\pi v_0}{L}\right)^2 \sin \frac{2\pi v_0 t}{L}$$

Le sinus est compris dans l'intervalle $[-1, 1]$, donc pour que l'accélération reste inférieure à $2g$ en valeur absolue, il faut :

$$a \left(\frac{2\pi v_0}{L}\right)^2 \leq 10g \Rightarrow a \leq 10g \left(\frac{L}{2\pi v_0}\right)^2$$

L'application numérique donne : $a = 9,9$ m.

Il faut être excellent pilote pour passer aussi près des cheminées à cette vitesse. Il faut supporter de plus une accélération importante ; à titre de comparaison lors du catapultage d'un avion depuis un porte-avion, l'accélération est environ $5g$. Mais dans le cas étudié, il s'agit de "Chevaliers Jedi" !

B.1 1. Pour le premier, on a une accélération a_1 constante soit après une double intégration l'équation horaire suivante $x_1 = \frac{1}{2}a_1 t^2$.

Pour le second, on a de même $x_2 = \frac{1}{2}a_2 (t - t_0)^2$ en tenant compte du retard t_0 au départ.

On cherche l'instant pour lequel on a $x_1 = x_2$. La résolution de l'équation du second degré $(a_1 - a_2)t^2 - 2a_2 t_0 t + a_2 t_0^2 = 0$ donne $t = \frac{a_2 \pm \sqrt{a_1 a_2}}{a_1 - a_2} t_0$.

L'application numérique donne $t = 9,5$ s et $t = 0,5$ s. La seconde solution n'est pas possible puisque l'un des deux n'est pas encore parti !

Il faut donc $8,5$ s pour que le second rattrape le premier.

2. Il suffit de reporter la valeur de t dans l'une ou l'autre des équations horaires. On trouve $x = 181$ m. On peut vérifier qu'on trouve bien la même valeur avec les deux équations horaires.
3. Quant aux vitesses, elles sont $v_1 = 38 \text{ m.s}^{-1}$ et $v_2 = 42,5 \text{ m.s}^{-1}$.

B.2 1. L'accélération est constante de D à E_1 , donc, à l'instant t , la vitesse est $v = a_1 t$ puisque la vitesse initiale est nulle et la distance parcourue est $d = a_1 t^2 / 2$, d'où :

$$t_{E_1} = \sqrt{\frac{2DE_1}{a_1}} = \sqrt{\frac{L}{a_1}} \quad \text{et} \quad v_{E_1} = a_1 t_{E_1} = \sqrt{a_1 L}$$

2. Dans le virage, le mouvement est circulaire, donc l'expression de l'accélération est :

$$\vec{a} = -R\dot{\theta}^2 \vec{u}_r + R\ddot{\theta} \vec{u}_\theta$$

D'après l'énoncé, on a $R\ddot{\theta} = a_1$, on en déduit en intégrant, $\dot{\theta} = a_1 t / R + cte$. On prend une nouvelle origine des temps en E_1 , alors $cte = \dot{\theta}_{E_1} = v_{E_1} / R = \sqrt{a_1 L} / R$. On calcule $\theta(t)$:

$$\theta = \frac{a_1 t^2}{2R} + \frac{\sqrt{a_1 L}}{R} t$$

en prenant $\theta = 0$ en E_1 . On obtient alors le temps de sortie du virage pour $\theta = \pi$, ce qui revient à résoudre l'équation du second degré :

$$t^2 + 2\sqrt{\frac{L}{a_1}} t - \frac{2\pi R}{a_1} = 0$$

La racine positive est :

$$t_1 = -\sqrt{\frac{L}{a_1}} + \sqrt{\frac{L + 2\pi R}{a_1}}$$

et la vitesse en S_1 est $v_{S_1} = R\dot{\theta}(S_1)$ soit :

$$v_{S_1} = a_1 t_1 + \sqrt{a_1 L} = \sqrt{(L + 2\pi R)a_1}$$

Le temps t_{S_1} demandé dans l'énoncé est à partir de l'origine en D soit :

$$t_{S_1} = t_{E_1} + t_1 = \sqrt{\frac{L + 2\pi R}{a_1}}$$

3. a) La distance entre S_1 et E_2 est égale à L . On change de nouveau l'origine des temps ($t = 0$ en S_1) et on note x l'abscisse à partir de S_1 . L'intégration de $\ddot{x} = a_1$ avec une vitesse initiale en S_1 donne : $x = \frac{a_1 t^2}{2} + v_{S_1} t$. On doit donc résoudre l'équation suivante pour calculer t_{E_2} :

$$t^2 + 2\frac{v_{S_1}}{a_1} t - \frac{2L}{a_1} = 0$$

La racine positive est :

$$t'_1 = -\sqrt{\frac{L + 2\pi R}{a_1}} + \sqrt{\frac{3L + 2\pi R}{a_1}}$$

d'où t_{E_2} :

$$t_{E_2} = t'_1 + t_{S_1} = \sqrt{\frac{3L + 2\pi R}{a_1}}$$

et $v_{E_2} = a_1 t'_1 + v_{S_1}$ soit :

$$v_{E_2} = \sqrt{(3L + 2\pi R)a_1}$$

En ce qui concerne t_{S_2} et v_{S_2} , le raisonnement est le même que pour le premier virage, seule la constante d'intégration pour θ change soit $\theta = \frac{a_1 t^2}{2R} + \frac{v_{E_2} t}{R}$, en prenant l'origine des temps et l'origine des angles en S_2 . L'équation donnant le temps de sortie en S_2 est (pour $\theta = \pi$) :

$$t^2 + 2\frac{v_{E_2}}{a_1} t - \frac{2R\pi}{a_1} = 0$$

La solution est :

$$t'_2 = -\sqrt{\frac{3L + 2\pi R}{a_1}} + \sqrt{\frac{3L + 4\pi R}{a_1}} \Rightarrow t_{S_2} = t_{E_2} + t'_2 = \sqrt{\frac{3L + 4\pi R}{a_1}}$$

et $v_{S_2} = \sqrt{(3L + 4\pi R)a_1}$

Enfin pour le calcul de t_D le raisonnement est le même que pour t_{E_2} , en changeant L en $L/2$ et v_{S_1} en v_{S_2} ce qui conduit à l'équation (en prenant les origines de x et t en S_2) :

$$t^2 + 2\frac{v_{S_2}}{a_1} t - \frac{L}{a_1} = 0$$

dont la solution positive est t'_3 :

$$t'_3 = -\sqrt{\frac{3L + 4\pi R}{a_1}} + 2\sqrt{\frac{L + \pi R}{a_1}} \Rightarrow t_D = t_{S_2} + t'_3 = 2\sqrt{\frac{L + \pi R}{a_1}}$$

et $v_D = 2\sqrt{(L + \pi R)a_1}$

- b) L'accélération a_1 est donnée par $a_1 = 4(L + \pi R)/t_D^2$ soit $a_1 = 1,515 \text{ m.s}^{-2}$. On obtient $v_D = 27,5 \text{ m.s}^{-1}$ ce qui fait 99 km.h^{-1} au lieu de 60 km.h^{-1} . On peut penser que l'accélération n'est pas constante et qu'elle diminue au cours du temps.

B.3 Étant donnée la symétrie du problème, les quatre robots sont toujours aux sommets d'un carré qui tourne en rétrécissant. Le point A_0 de la figure (S4.1) est le point de départ du robot A. L'axe OA_0 est la référence pour l'angle polaire θ et on note $r = OA$.

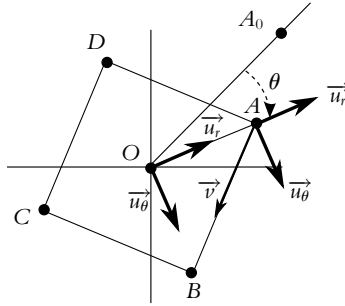


Figure S4.1

1. On remarque que $\vec{u}_r$ a la direction de la bissectrice de l'angle de sommet A donc la vitesse de A dirigée selon AB fait un angle $\pi/4$ avec $\vec{u}_\theta$ ainsi :

$$v_r = -v \cos \frac{\pi}{4} \text{ et } v_\theta = v \sin \frac{\pi}{4}$$

On connaît d'autre part, l'expression des composantes de la vitesse en coordonnées polaires d'où :

$$v_r = \dot{r} = -\frac{v}{\sqrt{2}} \text{ et } v_\theta = r\dot{\theta} = \frac{v}{\sqrt{2}}$$

On intègre l'équation pour r : $r = -\frac{vt}{\sqrt{2}} + cte$ or pour $t = 0$, $r = \frac{a}{\sqrt{2}}$, d'où :

$$r(t) = -\frac{vt}{\sqrt{2}} + \frac{a}{\sqrt{2}}$$

L'équation en θ devient $\dot{\theta} = \frac{v}{a - vt}$, ce qui donne par intégration avec $\theta(t = 0) = 0$:

$$\theta(t) = \ln \frac{a}{a - vt}$$

2. L'instant d'arrivée en O correspond à $r = 0$ soit $t = a/v$. On remarque que $\theta \rightarrow \infty$, ce qui signifie que les robots arrivent au centre au bout d'un temps fini, mais après avoir effectué un nombre infini de tours. Ce paradoxe est dû au fait que l'on a considéré que les robots étaient des points matériels sans dimension.

B.4 1. Le point a un mouvement circulaire uniforme de centre I. Il décrit un tour en une période et un demi-tour de O à C en une demi-période soit $T = 2\pi/\omega$.

2. a) La relation entre φ et θ est une propriété des cercles : $\varphi = 2\theta$. On peut la démontrer en considérant le triangle isocèle OIB. Si l'on note β les angles (IOB) et (IBO), alors dans le triangle OIB, $2\beta + \varphi = \pi$, or $\theta = \pi/2 - \beta$, d'où le résultat.

- b) L'angle φ correspond à l'angle polaire du mouvement circulaire de centre I donc $\dot{\varphi} = \omega$. Puisque ω est constant, on a $\varphi = \omega t + cte$, or à $t = 0$, B est en O donc $\varphi(0) = 0$ et $cte = 0$. On en déduit $\theta = \omega t/2$.

En ce qui concerne $r = OB$, dans le triangle OIB , on peut écrire $OB = 2b \sin \frac{\varphi}{2}$, soit $r = 2b \sin \frac{\omega t}{2}$.

3. a) Les coordonnées de A dans la base $(\vec{u}_r, \vec{u}_\theta, \vec{u}_z)$ sont $(0, 0, 2b \cos \alpha)$; celles de B sont $(r, 0, 0)$ d'où $\vec{AB} = r\vec{u}_r - 2b \cos \alpha \vec{u}_z$. Or la norme de $\vec{AB}$ est $2b$. On en déduit :

$$4b^2 \sin^2 \frac{\omega t}{2} + 4b^2 \cos^2 \alpha = 4b^2 \Rightarrow \cos^2 \alpha = 1 - \sin^2 \frac{\omega t}{2}$$

comme $\alpha \in [0, \pi/2]$, alors $\alpha = \frac{\omega t}{2}$.

- b) Initialement la barre est suivant l'axe Oz et dans l'état final suivant l'axe Oy .

4. Coordonnées de J :

$$\begin{cases} X = \frac{1}{2}OB \cos \theta = \frac{1}{2}r \cos \theta = b \sin \frac{\omega t}{2} \cos \frac{\omega t}{2} = \frac{1}{2}b \sin \omega t \\ Y = \frac{1}{2}OB \sin \theta = b \sin^2 \frac{\omega t}{2} \\ Z = \frac{1}{2}2b \cos \alpha = b \cos \frac{\omega t}{2} \end{cases}$$

5. On dérive les coordonnées de J soit :

$$\vec{v} \rightarrow \begin{cases} \dot{X} = \frac{\omega}{2}b \cos \omega t \\ \dot{Y} = \omega b \sin \frac{\omega t}{2} \cos \frac{\omega t}{2} = \frac{\omega}{2}b \sin \omega t \\ \dot{Z} = -\frac{\omega}{2}b \sin \frac{\omega t}{2} \end{cases}$$

et pour l'accélération :

$$\vec{a} \rightarrow \begin{cases} \ddot{X} = -\frac{\omega^2}{2}b \sin \omega t \\ \ddot{Y} = \frac{\omega^2}{2}b \cos \omega t \\ \ddot{Z} = -\frac{\omega^2}{4}b \cos \frac{\omega t}{2} \end{cases}$$

On en déduit les normes :

$$\|\vec{v}\| = b \frac{\omega}{2} \sqrt{1 + \sin^2 \frac{\omega t}{2}} \quad \text{et} \quad \|\vec{a}\| = b \frac{\omega^2}{2} \sqrt{1 + \frac{1}{4} \cos^2 \frac{\omega t}{2}}$$

Chapitre 5

A.1 Dans le référentiel du laboratoire supposé galiléen, les forces s'exerçant sur le système M sont : le poids $m\vec{g}$ et la tension du fil $\vec{T}$ (figure S5.1). On utilise comme base de projection, la base cylindrique $(\vec{u}_r, \vec{u}_\theta, \vec{u}_z)$ et les coordonnées correspondantes : $r = HM$, et θ tel que $\dot{\theta} = \omega$.

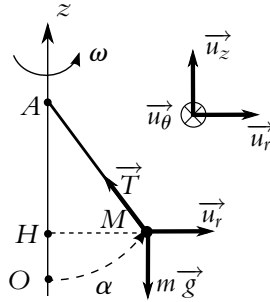


Figure S5.1

On applique la relation fondamentale de la dynamique à M : $m\vec{a} = m\vec{g} + \vec{T}$. On note $T = \|\vec{T}\|$ et $g = \|\vec{g}\|$ soit en projection :

$$\begin{cases} m(\ddot{r} - r\dot{\theta}^2) = -T \sin \alpha \\ m(2\dot{r}\dot{\theta} + r\ddot{\theta}) = 0 \\ m\ddot{z} = -mg + T \cos \alpha \end{cases}$$

Sachant que ω est constante, $\ddot{\theta} = 0$ et la deuxième équation devient :

$$2m\dot{r}\omega = 0 \Rightarrow \dot{r} = 0$$

Dans le triangle AMH , on peut écrire : $l^2 = r^2 + (l - z)^2$, donc puisque r est constant, z l'est aussi ainsi que α . En remplaçant r par $l \sin \alpha$, on obtient : $ml\omega^2 \sin \alpha = T \sin \alpha$ et $T \cos \alpha = mg$. Deux solutions sont possibles :

$$\begin{cases} \alpha = 0 \text{ et } T = mg \\ T = ml\omega^2 \text{ et } \cos \alpha = \frac{g}{l\omega^2} \end{cases}$$

La première solution n'est pas stable ce qui pourra être montré lorsque les connaissances sur les équilibres dans les référentiels non galiléens seront acquises. Ainsi seule la seconde solution est observable physiquement.

A.2 On suppose le référentiel lié à la Lune galiléen. Le système est le sauteur qu'on nomme M . Le repère de projection est (O, x, y, z) , le point O est l'origine du saut et Oz est la verticale ascendante. On a choisis les axes tels que

$$\vec{v}_0 = v_0(\cos \alpha \vec{u}_x + \sin \alpha \vec{u}_z)$$

La seule force est $m\vec{g}_l = -mg_l\vec{u}_z$.

1. La relation fondamentale de la dynamique donne : $m\vec{a} = m\vec{g}_l$, soit par intégration successive

$$\begin{cases} \ddot{x} = 0 \\ \ddot{y} = 0 \\ \ddot{z} = -g_l \end{cases} \Rightarrow \begin{cases} \dot{x} = v_0 \cos \alpha \\ \dot{y} = 0 \\ \dot{z} = -g_l t + v_0 \sin \alpha \end{cases} \Rightarrow \begin{cases} x = v_0 t \cos \alpha \\ y = 0 \\ z = -\frac{1}{2}g_l t^2 + v_0 t \sin \alpha \end{cases}$$

On obtient l'équation de la trajectoire en éliminant t :

$$z = -\frac{1}{2}g_l \frac{x^2}{(v_0 \cos \alpha)^2} + x \tan \alpha$$

Le sauteur retombe sur le sol pour $z = 0$, soit en rejetant la solution $x = 0$ qui correspond au départ, on trouve la distance parcourue :

$$d = \frac{v_0^2}{g_l} \sin 2\alpha$$

2. Puisque d est inversement proportionnel à g_l , la distance sur la Lune sera six fois plus grande soit 9 m.

A.3 Le référentiel lié au sol est considéré galiléen. Le système est la masse M . On nomme $\vec{u}_z$ le vecteur unitaire de l'axe Oz . On attribuera l'indice (1) aux grandeurs relatives au ressort supérieur, et l'indice (2) au ressort inférieur.

On a défini pour chaque ressort le vecteur unitaire $\vec{u}_{\text{ext}}$ orienté du point d'attache vers le point mobile. Si la longueur du ressort est ℓ , la tension exercée par le ressort sur M est $\vec{T} = -k(\ell - \ell_0)\vec{u}_{\text{ext}}$.

Dans le cas présent, $\vec{u}_{\text{ext},1} = -\vec{u}_z$ et $\vec{u}_{\text{ext},2} = \vec{u}_z$.

Les trois forces s'exerçant sur M sont le poids, la tension $\vec{T}_1$ et la tension $\vec{T}_2$.

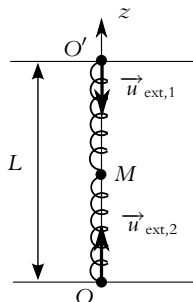


Figure S5.2

Les tensions s'écrivent, $\vec{T}_1 = -k(\ell_1 - \ell_0)\vec{u}_{\text{ext},1}$ et $\vec{T}_2 = -k(\ell_2 - \ell_0)\vec{u}_{\text{ext},2}$. L'origine est prise en O donc $\ell_1 = L - z$ et $\ell_2 = z$.

1. La relation fondamentale à l'équilibre s'écrit :

$$m\vec{g} + \vec{T}_1 + \vec{T}_2 = \vec{0} \Rightarrow -mg + k(L - z_e - \ell_0) - k(z_e - \ell_0) = 0 \Rightarrow z_e = \frac{L}{2} - \frac{mg}{2k}$$

En l'absence de poids, la position d'équilibre serait au milieu, alors que dans le cas présent elle est en-dessous.

2. On écrit la relation fondamentale de la dynamique :

$$m\vec{a} = m\vec{g} + \vec{T}_1 + \vec{T}_2 \Rightarrow m\ddot{z} = -mg + k(L - z - \ell_0) - k(z - \ell_0)$$

En soustrayant de cette équation celle à l'équilibre il vient :

$$m\ddot{z} = -2kz + 2kz_e \Rightarrow \ddot{z} + \omega_0^2 z = \omega_0^2 z_e$$

3. La solution générale de cette équation est : $z = A \cos \omega_0 t + B \sin \omega_0 t + z_e$. À $t = 0$, $z = z_e + a$ et $\dot{z} = 0$. Avec $\dot{z} = \omega_0(-A \sin \omega_0 t + B \cos \omega_0 t)$, on obtient à $t = 0$:

$$A = a \text{ et } \omega_0 B = 0$$

d'où $z(t) = a \cos \omega_0 t$.

A.4 Le référentiel lié au sol est galiléen. Le système est le point M .

1. Lorsque M est posé sur le sol, les forces s'exerçant sur le point sont : le poids ($m\vec{g}$), la tension du câble ($\vec{T}$) et la réaction du sol ($\vec{R}$). À l'équilibre : $m\vec{g} + \vec{T} + \vec{R} = \vec{0}$. La condition pour que la masse se soulève juste à la rupture est donnée par $\vec{R} = \vec{0}$, d'où $\vec{T} = -m\vec{g}$.

2. On applique le relation fondamentale de la dynamique avec $\vec{a} = \vec{a}_v$, soit $m\vec{a}_0 = m\vec{g} + \vec{T}$, d'où $\vec{T} = m(\vec{a}_v - \vec{g})$, soit en norme $T = m(a_v + g)$, $\vec{a}_v$ étant vers le haut.

La tension est supérieure au poids et fonction affine de a_0 . Si l'accélération était trop forte, le câble pourrait se rompre.

3. a) Si l'on note O un point fixe dans le référentiel, alors l'accélération de M est $\vec{a}_M = \frac{d\vec{OM}}{dt}$. On utilise la relation de Chasles $\vec{OM} = \vec{OA} + \vec{AM}$, or $\vec{AM}$ est un vecteur constant, donc :

$$\vec{a}_M = \frac{d\vec{OM}}{dt} = \frac{d\vec{OA}}{dt} + \frac{d\vec{AM}}{dt} = \frac{d\vec{OA}}{dt} = \vec{a}_h$$

- b) La relation fondamentale de la dynamique appliquée à M donne $m\vec{a}_h = m\vec{g} + \vec{T}$, les différents vecteurs étant représentés sur la figure (S5.3). La construction vectorielle donne alors $\tan \alpha = mg/a_h$ et $T = \sqrt{(mg)^2 + a_h^2}$.

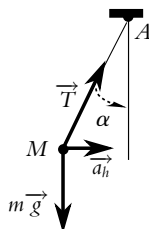


Figure S5.3

B.1 1. Le système étudié est le ballon dans le référentiel terrestre galiléen. Il est soumis à la seule force du poids. L'écriture du principe fondamental de la dynamique donne $\vec{a} = \vec{g}$.

2. La projection de la relation précédente en coordonnées cartésiennes en choisissant Oz dirigé vers le haut donne $\ddot{x} = 0$, $\ddot{y} = 0$, $\ddot{z} = -g$. La vitesse s'obtient par intégration, ce qui donne compte tenu des conditions initiales $\dot{x} = v_0 \cos \alpha$, $\dot{y} = 0$, $\dot{z} = -gt + v_0 \sin \alpha$ puis une nouvelle intégration permet d'obtenir la position soit $x = v_0 t \cos \alpha - D$, $y = 0$, $z = \frac{1}{2}gt^2 + v_0 t \sin \alpha + h$ en prenant l'origine au sol sous le panneau.

3. Pour obtenir l'équation de la trajectoire, il suffit d'éliminer le temps entre les différentes équations. On en déduit $t = \frac{x + D}{v_0 \cos \alpha}$ donc

$$z = \frac{1}{2}g \left(\frac{x + D}{v_0 \cos \alpha} \right)^2 + (x + D) \tan \alpha + h$$

4. La résolution de l'équation précédente pour $z = H$ et $x = 0$ donne, en transformant $\frac{1}{\cos^2 \alpha} = 1 + \tan^2 \alpha$, $\frac{gD^2}{2v_0^2} \tan^2 \alpha - D \tan \alpha - h + \frac{gD^2}{2v_0^2} + H = 0$.

5. Les solutions doivent être réelles pour avoir un sens physiquement. Ce sera uniquement le cas si le discriminant est positif soit $D^2 - 4 \frac{gD^2}{2v_0^2} \left(-h + \frac{gD^2}{2v_0^2} + H \right) > 0$. Cette condition peut s'écrire $v_0^4 - 2g(H - h)v_0^2 - g^2D^2 \geq 0$.

6. Le discriminant de ce polynôme bicarré est $\delta = g^2(H - h)^2 + g^2D^2 > 0$. Les deux racines du polynôme sont $v_0^2 = g(H - h) \pm g\sqrt{(H - h)^2 + D^2}$, l'une positive et l'autre négative. Par conséquent, l'expression est donc entre les racines du signe de la valeur prise pour $v_0^2 = 0$ soit négative. Comme $v_0^2 \geq 0$, la condition est vérifiée si v_0^2 est supérieure à la racine positive soit

$$v_0 \geq \sqrt{g \left(H - h + \sqrt{(H - h)^2 + D^2} \right)}.$$

7. L'application numérique donne pour un lancer franc

$$v_0 \geq 7,59 \text{ m.s}^{-1} = 27,3 \text{ km.h}^{-1}$$

et pour un tir à trois points

$$v_0 \geq 8,60 \text{ m.s}^{-1} = 31,0 \text{ km.h}^{-1}.$$

8. La résolution de l'équation du second degré en $\tan \alpha$ conduit à des solutions réelles puisque les conditions précédentes sont vérifiées. Elles s'écrivent $\tan \alpha = \frac{v_0^2 \pm \sqrt{v_0^4 - 2g(H-h)v_0^2 - g^2 D^2}}{gD}$.

On a donc deux valeurs possibles pour l'angle α .

9. L'application numérique donne $\tan \alpha = 0,52$ ou $3,93$ soit $\alpha = 27^\circ 28'$ ou $\alpha = 75^\circ 42'$.

10. La résolution à α fixé de l'équation en v_0 donne

$$v_0 = \sqrt{\frac{gD^2}{2 \cos^2 \alpha (h - H + D \tan \alpha)}}.$$

11. L'application numérique donne $v_0 = 8,75 \text{ m.s}^{-1}$.

12. Le bord du ballon se situe entre 0,00 cm et 45 cm donc son centre est entre 17,8 cm et 27,2 cm.

13. D'après la question 3, en posant $X = x + D$, on doit résoudre

$$gX^2 - Xv_0^2 \sin 2\alpha + 2v_0^2 (H-h) \cos^2 \alpha = 0$$

14. Cette équation admet au moins une solution si son discriminant est positif soit $v_0^2 \sin^2 \alpha \geq 2g(H-h)$.

15. A v_0 fixée, il faut donc que $\sin^2 \alpha \geq 0,21$ et comme on effectue le tir en l'air et devant, l'angle α vérifie $0 \leq \alpha \leq 90^\circ$. On en déduit finalement que $26^\circ 59' \leq \alpha \leq 90^\circ$.

16. A α fixé, la condition s'écrit $v_0 \geq 4,83 \text{ m.s}^{-1}$.

17. On écrit la solution $x = \frac{v_0^2 \sin 2\alpha \pm v_0 \sqrt{v_0^2 \sin^2 2\alpha - 8(H-h) \cos^2 \alpha}}{2g} - D$.

B.2 1. Le mobile démarre avec une vitesse nulle, la force de frottement est donc nulle aussi. La force qui le fait descendre est le poids qui est une force constante, la vitesse augmente, donc la force de frottement (proportionnelle à v^2) augmente aussi. Lorsque la force de frottement devient égale au poids, la force totale s'exerçant sur la mobile est alors nulle et la vitesse constante : c'est la vitesse limite. On a donc :

$$mg = kv_l^2 \Rightarrow v_l = \sqrt{\frac{mg}{k}}$$

2. a) On écrit la relation fondamentale que l'on projette sur l'axe Oz descendant :

$$m \vec{a} = m \vec{g} - kv^2 \vec{u}_z \Rightarrow m \ddot{z} = mg - k\dot{z}^2 = mg \left(1 - \left(\frac{\dot{z}}{v_l} \right)^2 \right)$$

soit avec $v = \dot{z}$:

$$\frac{\dot{v}}{1 - \left(\frac{v}{v_l} \right)^2} = g$$

Il s'agit d'une équation différentielle non linéaire.

- b) Sachant que $\dot{v} = \frac{dv}{dt}$:

$$\int \frac{dv}{1 - \left(\frac{v}{v_l} \right)^2} = \int g dt \Rightarrow v_l \tanh^{-1} \left(\frac{v}{v_l} \right) = gt + \text{cte}$$

Or à $t = 0$, $v = 0$ soit cte=0. On en déduit :

$$v(t) = v_l \tanh \frac{gt}{v_l} = v_l \tanh \frac{t}{\tau}$$

Sachant que $\tanh(3) = 0,995$, le mobile aura une vitesse égale à $0,995v_l$ au bout de $t = 3\tau$.

- c) Pour calculer $z(t)$, il faut intégrer $v = \frac{dz}{dt} = v_l \tanh \frac{t}{\tau}$, soit :

$$z(t) = \int v_l \tanh \left(\frac{t}{\tau} \right) dt + cte = v_l \int \frac{\sinh(t/\tau)}{\cosh(t/\tau)} dt + cte = v_l \tau \ln \left(\cosh \left(\frac{t}{\tau} \right) \right) + cte$$

À $t = 0$, $z = 0$ donc $cte = 0$.

La cote $z(t)$ dépend de m puisque τ dépend de v_l qui dépend lui-même de m . S'il n'y avait pas de frottement, la seule force serait le poids et la masse m se simplifierait dans l'équation du mouvement. Dans le vide, tout corps suit la même loi de chute libre (expérience du tube de Newton avec une plume et une masse).

3. On rappelle l'expression de $v_l = \sqrt{\frac{mg}{k}}$ et $\tau = \frac{v_l}{g}$.

a) Vitesses limites : pour la balle de tennis $v_{l,1} = 16,2 \text{ m.s}^{-1}$ et pour la boule de pétanque $v_{l,2} = 49,8 \text{ m.s}^{-1}$, donc $\tau_1 = 1,78 \text{ s}$ et $\tau_2 = 5,08 \text{ s}$. La différence est essentiellement due à la différence de masse.

b) On obtient le temps de chute en inversant l'expression de $z(t)$:

$$t = \tau \cosh^{-1} \left(\exp \left(\frac{z}{v_l \tau} \right) \right)$$

L'application numérique donne $t_1 = 635 \text{ ms}$, $t_2 = 607 \text{ ms}$ et $\Delta t = 28,8 \text{ ms}$. Les corps n'ont pas eu le temps d'atteindre leur vitesse limite. Lorsque la boule de pétanque arrive au sol, la cote de la balle de tennis est $z_1(0,607 \text{ s}) = 1,645 \text{ m}$ soit $15,5 \text{ cm}$ ce qui est notable.

B.3 1. Le référentiel lié au sol est galiléen. On définit le repère A, AX, AY comme sur la figure (S5.4). Les forces s'exerçant sur P sont le poids, la réaction normale et la réaction tangentielle.

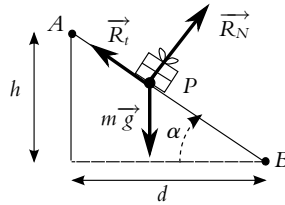


Figure S5.4

On écrit la relation fondamentale de la dynamique : $m \vec{a} = m \vec{g} + \vec{R}_t + \vec{R}_N$. La force $\vec{R}_t$ est une force de frottement donc elle s'oppose au mouvement. Les forces sont représentées sur la figure (S5.4). Les projections sur AX et AY donnent :

$$\begin{cases} m\ddot{X} = mg \sin \alpha - R_t \\ m\ddot{Y} = -mg \cos \alpha + R_N \end{cases}$$

où g , R_t et R_N sont les normes des vecteurs correspondants. Or $Y = cte$ donc $\ddot{Y} = 0$ soit $R_N = mg \cos \alpha$. On en déduit $R_t = fmg \cos \alpha$. On reporte dans l'équation sur AX et on intègre :

$$\ddot{X} = g(\sin \alpha - f \cos \alpha) \Rightarrow \dot{X} = gt(\sin \alpha - f \cos \alpha) + cte$$

La constante est nulle car la vitesse initiale est nulle ; puis :

$$X = \frac{gt^2}{2} (\sin \alpha - f \cos \alpha) + cte$$

La constante est nulle car $X(0) = 0$.

La position du point B correspond à $X = \frac{h}{\sin \alpha}$, donc B est atteint en un temps t_B tel que :

$$h = \frac{gt_B^2}{2} (\sin \alpha - f \cos \alpha) \sin \alpha \Rightarrow t_B = \sqrt{\frac{2h}{g \sin \alpha (\sin \alpha - f \cos \alpha)}}$$

Puisque $h = d$, l'angle α vaut $\pi/4$ soit $t_B = 1,8 \text{ s}$.

2. Le chariot reste immobile $\delta t = 1$ s, il faut donc que le paquet parte au plus tard à $t_2 = t_B - \delta t = 0,8$ s avant l'arrivée du chariot et au plus tôt $t_B = 1,8$ s avant l'arrivée.

On calcule l'instant t_a d'arrivée du chariot en C_B : $t_a = C_0 C_B / v_c = 4$ s. Le joueur doit donc lâcher le paquet entre 2,2 s et 3,2 s.

B.4 1. Attention, ce n'est pas le poids du sauteur qui est donné, mais sa masse.

La tension est une force et s'exprime donc en N et non en kg. Ce qui est donné dans la fiche, correspond vraisemblablement à la masse m qui, accrochée à l'élastique donne l'allongement correspondant. Par application de la relation fondamentale à l'équilibre, on obtient $T = mg$, soit $T_1 = 1960$ N pour un allongement de 100% et $T_2 = 3185$ N pour un allongement de 200%.

2. a) Par définition l'allongement de l'élastique (en %) est $\frac{\ell - \ell_0}{\ell_0}$. Un allongement de 100% correspond à $\ell = 2\ell_0$ et de 200% à $\ell = 3\ell_0$. Ainsi dans le deuxième cas, si la tension suivait la loi proposée, elle devrait être égale à 1,5 fois la première soit à 2940 N. Elle est en fait plus grande, ce qui signifie que lorsque l'élastique est très étiré k n'est plus une constante ; sa valeur augmente.
- b) Pour l'allongement de 100%, on trouve $k_1 = T_1 / 2\ell_0 = 163 \text{ N.m}^{-1}$ et pour l'allongement de 200%, on trouve $k_2 = T_2 / 3\ell_0 = 180 \text{ N.m}^{-1}$ soit une valeur moyenne $k = 171,5 \text{ N.m}^{-1}$ qui sera celle retenue dans la suite.
3. a) Le système est le point M de masse m qui est soumis à son poids $m\vec{g}$ et à la tension de l'élastique $\vec{T}$ et la relation fondamentale donne $m\vec{a} = m\vec{g} + \vec{T}$. Le problème qui se pose avec les forces de tension d'élastique (ou de ressort) est de bien les orienter lors de la projection ; une méthode efficace est de définir le vecteur $\vec{u}_{\text{ext}}$; ce vecteur s'oriente du point d'attache de l'élastique (ou du ressort) vers la masse en mouvement : on peut alors écrire $\vec{T} = -k(\ell - \ell_0)\vec{u}_{\text{ext}}$. La projection de la relation fondamentale sur un axe vertical Oz descendant ($\vec{u}_z = \vec{u}_{\text{ext}}$) avec O point d'attache du ressort (alors $\ell = z$) donne alors :

$$m\ddot{z} = mg - k(\ell - \ell_0) \Rightarrow \ddot{z} + \frac{k}{m}z = g + \frac{k}{m}\ell_0$$

La solution de l'équation est $z = A \cos(\omega_0 t) + B \sin(\omega_0 t) + \ell_0 - \frac{mg}{k}$, avec $\omega_0 = \sqrt{\frac{k}{m}}$. Dans la suite, on pose $\ell_c = \ell_0 + \frac{mg}{k}$.

On peut alors calculer la vitesse : $\dot{z} = \omega_0(-A \sin(\omega_0 t) + B \cos(\omega_0 t))$. On détermine les constantes A et B avec les conditions initiales, à $t = 0$:

$$\begin{cases} z(0) = h \\ \dot{z}(0) = 0 \end{cases} \Rightarrow \begin{cases} A + \ell_c = h \\ B\omega_0 = 0 \end{cases}$$

d'où $z(t) = (h - \ell_c) \cos(\omega_0 t) + \ell_c$ et $\dot{z} = -\omega_0(h - \ell_c) \sin(\omega_0 t)$.

- b) L'élastique se détend à l'instant t_1 tel que $z(t) = \ell_0$ soit :

$$t_1 = \frac{1}{\omega_0} \text{Arccos} \left[-\frac{mg}{k} \left(h - \frac{mg}{k} - \ell_0 \right)^{-1} \right] = 1,25 \text{ s}$$

La vitesse $\overline{v_1}$ à l'instant t_1 est alors $\dot{z}(t_1) = -3,7 \text{ m.s}^{-1}$.

- c) Pour $\ell < \ell_0$, M n'est plus soumise qu'à son poids :

$$\ddot{z} = g \Rightarrow \dot{z} = g(t - t_1) + \overline{v_1} \Rightarrow z(t) = \frac{g}{2}(t - t_1)^2 + \overline{v_1}(t - t_1) + \ell_0$$

Il faut prendre garde au fait que les conditions initiales pour cette étape du mouvement doivent être prises à l'instant t_1 et que l'axe est orienté vers le bas, d'où la projection de $\vec{g}$. La hauteur maximale atteinte, l'est lorsque la vitesse est nulle soit $t - t_1 = -\overline{v_1}/g$ et donc :

$$z_{\text{max}} = -\frac{g}{2} \left(\frac{\overline{v_1}}{g} \right)^2 - \overline{v_1} \frac{\overline{v_1}}{g} + \ell_0 = \ell_0 - \frac{3\overline{v_1}^2}{2g} = 3,9 \text{ m}$$

4. a) Le saut se passe en deux étapes (inversées par rapport à la précédente partie) : une chute libre jusqu'à $z = \ell_0$ puis une chute soumise au poids et à la tension de l'élastique. Il faut donc calculer

la vitesse atteinte à la fin de la première étape en ℓ_0 pour avoir les conditions initiales pour la deuxième étape.

Pour la première étape, le départ ayant lieu pour $z = 0$ au point d'attache de l'élastique (axe Oz vertical descendant comme précédemment) avec vitesse initiale nulle :

$$\ddot{z} = g \Rightarrow \dot{z} = gt \Rightarrow z(t) = \frac{g}{2}t^2$$

donc pour $z = \ell_0$, $t_2 = \sqrt{\frac{2\ell_0}{g}}$ et $\dot{z}(\ell_0) = \sqrt{2g\ell_0}$.

Pour la deuxième étape, la relation fondamentale donne :

$$m\vec{a} = m\vec{g} - k(z - \ell_0)\vec{u}_z.$$

Cette équation a déjà été résolue précédemment, cependant comme cette fois-ci on nous demande la hauteur maximale de chute, il est préférable de prendre la solution sous la forme : $z(t) = A' \cos(\omega_0 t + \varphi) + \ell_0 + \frac{mg}{k}$ (où $\omega_0 = \sqrt{k/m}$). Avec les conditions initiales à $t = t_2$, on détermine A' (pas besoin de déterminer φ) :

$$\begin{cases} z(t_2) = A' \cos(\omega_0 t_2 + \varphi) + \ell_0 + \frac{mg}{k} = \ell_0 \\ \dot{z}(t_2) = -A' \omega_0 \sin(\omega_0 t_2 + \varphi) = \sqrt{2g\ell_0} \end{cases} \Rightarrow A' = \sqrt{\left(\frac{mg}{k}\right)^2 + \frac{2mg\ell_0}{k}}$$

Pour obtenir A' , on a élevé chaque équation au carré et utilisé la relation $\cos^2 + \sin^2 = 1$. L'application numérique donne : $A' = 7,6$ m. La hauteur maximale de chute est obtenue pour $\cos(\omega_0 t_2 + \varphi) = 1$, soit $z_{\max} = A' + \ell_0 + \frac{mg}{k}$ soit $z_{\max} = 17,3$ m ce qui est pratiquement la hauteur du pont.

L'accélération est obtenue par deux dérivations de $z(t)$ soit : $\ddot{z} = -A' \omega_0^2 \cos(\omega_0 t_2 + \varphi)$ soit $|\ddot{z}|_{\max} = A' \omega_0^2 = 20,1 \text{ m.s}^{-2}$, soit plus de deux fois l'accélération de la pesanteur.

- b)** Dans l'équation du mouvement, la position d'équilibre est obtenue pour $\ddot{z} = 0$, soit $z = \ell_e = \ell_0 + \frac{mg}{k}$. C'est autour de cette position qu'ont lieu les oscillations du sauteur.

B.5 On considère le référentiel terrestre galiléen. Le système est le plongeur.

- 1.** La première partie du mouvement correspond à une chute libre. La relation fondamentale de la dynamique $m\vec{a} = m\vec{g}$ donne en projection sur Oz descendant $\ddot{z} = g$, soit en intégrant avec $\dot{z}(0) = 0$ et $z(0) = 0$:

$$\dot{z} = gt \text{ et } z(t) = \frac{gt^2}{2}$$

Pour $z = h$, on trouve $t_c = \sqrt{2h/g} = 1,4$ s et $v_c = \sqrt{2gh} = 14,1 \text{ m.s}^{-1}$.

- 2. a)** La relation fondamentale dans l'eau devient : $m\vec{a} = m\vec{g} + \vec{\Pi} + \vec{f}_f$, ce qui donne en projection sur Oz :

$$m\dot{v}_z = mg \left(1 - \frac{1}{d_h}\right) - kv_z \Rightarrow \dot{v}_z + \frac{v_z}{\tau} = g \left(1 - \frac{1}{d_h}\right)$$

- b)** La solution générale de l'équation précédente est :

$$v_z = A \exp\left(-\frac{t}{\tau}\right) + g\tau \left(1 - \frac{1}{d_h}\right)$$

soit en changeant l'origine des temps, $v(t=0) = v_e$:

$$v_z(t) = g\tau \left(1 - \frac{1}{d_h}\right) + \left(v_e - g\tau \left(1 - \frac{1}{d_h}\right)\right) \exp\left(-\frac{t}{\tau}\right)$$

- c)** Lorsque $t \rightarrow \infty$, on trouve la vitesse limite, soit $v_L = g\tau \left(1 - \frac{1}{d_h}\right)$. L'application numérique donne $v_L = -0,356 \text{ m.s}^{-1}$.

- d) On peut alors mettre v_z sous la forme suivante :

$$v_z(t) = (v_e + |v_L|) \exp\left(-\frac{t}{\tau}\right) - |v_L|$$

La baigneur commence à remonter lorsque la vitesse v_z devient nulle, puisque la résultante des forces est vers le haut :

$$|v_L| = (v_e + |v_L|) \exp\left(-\frac{t_1}{\tau}\right) \Rightarrow t_1 = \tau \ln\left(1 + \frac{v_e}{|v_L|}\right) = 1,19 \text{ s}$$

- e) En intégrant la vitesse, avec $z(t=0) = 0$, on trouve :

$$z(t) = \tau(v_e + |v_L|) \left(1 - \exp\left(-\frac{t}{\tau}\right)\right) - |v_L|t$$

La profondeur maximale atteinte l'est pour $t = t_1$, soit 4,1 m.

- f) On cherche l'instant t_2 pour lequel $v_z = v_2 = 1 \text{ m.s}^{-1}$. la résolution donne :

$$t_2 = \tau \ln\left(\frac{v_e + |v_L|}{v_2 + |v_L|}\right) = 0,76 \text{ s}$$

Alors $z = 3,94 \text{ m}$.

3. a) Les forces sont les mêmes que précédemment. La seule force qui a une projection selon l'horizontale est la force de frottement, ce qui donne :

$$m\ddot{x} = -k\dot{x} \quad \text{et} \quad m\ddot{z} = mg\left(1 - \frac{1}{d_h}\right) - k\dot{z}$$

- b) L'intégration des équations donne :

$$\dot{x} = A \exp\left(-\frac{t}{\tau}\right) \quad \text{et} \quad \dot{z} = B \exp\left(-\frac{t}{\tau}\right) + g\tau\left(1 - \frac{1}{d_h}\right)$$

Sachant qu'à $t = 0$ (instant de pénétration dans l'eau) $\dot{x}(0) = v_0 \cos \alpha$ et $\dot{z}(0) = v_0 \sin \alpha$, on obtient :

$$\dot{x} = v_0 \cos \alpha \exp\left(-\frac{t}{\tau}\right)$$

et

$$\dot{z} = g\tau\left(1 - \frac{1}{d_h}\right) + \left(v_0 \sin \alpha - g\tau\left(1 - \frac{1}{d_h}\right)\right) \exp\left(-\frac{t}{\tau}\right)$$

La vitesse limite est obtenue pour $t \rightarrow \infty$. On trouve une vitesse limite verticale d'expression semblable à celle calculée précédemment. On peut mettre z sous une forme semblable à celle de la question précédente en remplaçant v_e par $v_0 \sin \alpha$.

- c) Toutes les expressions obtenues pour le saut, sont applicables dans le cas du plongeon en remplaçant v_e par $v_0 \sin \alpha$. Avec les nouvelles valeurs de k et donc de $\tau = 0,64 \text{ s}$ et de $v_L = -0,71 \text{ m.s}^{-1}$, la vitesse s'annule donc à l'instant :

$$t_3 = \tau \ln\left(1 + \frac{v_0 \sin \alpha}{|v_L|}\right) = 1,52 \text{ s}$$

et à la profondeur :

$$z(t_3) = \tau(v_0 \sin \alpha + |v_L|) \left(1 - \exp\left(-\frac{t_3}{\tau}\right)\right) - |v_L|t_3 = 3,35 \text{ m}$$

Le plongeur n'atteint donc pas le fond.

- B.6** 1. Dans le référentiel lié à la route considéré galiléen, les forces s'exerçant sur le système M sont :

- Le poids $m\vec{g}$;
- la réaction normale du support $\vec{R}_N$;
- la force motrice $\vec{F}$.

La relation fondamentale donne : $m \vec{a} = m \vec{g} + \vec{R}_N + \vec{F}$. La force motrice est tangente à la trajectoire, donc $\vec{F} = \overline{F} \vec{u}_\theta$.

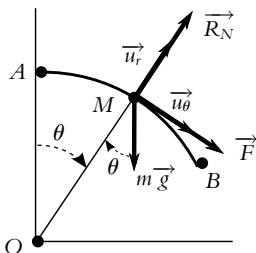


Figure S5.5

En projection sur la base polaire, on obtient ($r = R = \text{cte}$) :

$$\begin{cases} m(\ddot{r} - r\dot{\theta}^2) = -mR\dot{\theta}^2 = R_N - mg \cos \theta \\ m(2\dot{r}\dot{\theta} + r\ddot{\theta}) = mR\ddot{\theta} = mg \sin \theta + \overline{F} \end{cases}$$

2. On multiplie l'équation sur $\vec{u}_\theta$ par $\dot{\theta}$ soit $mR\ddot{\theta}\dot{\theta} = mg\dot{\theta} \sin \theta + \overline{F}\dot{\theta}$, dont l'intégration donne :

$$mR \frac{\dot{\theta}^2}{2} = -mg \cos \theta + \overline{F}\theta + \text{cte}$$

On remarque que pour intégrer l'équation par rapport au temps, il est nécessaire de la multiplier par $\dot{\theta}$, car $\frac{d(\cos \theta)}{dt} = -\dot{\theta} \sin \theta$.

Les conditions initiales sont, pour $\theta = 0$, $\dot{\theta} = v_0/R$ soit $\text{cte} = mv_0^2/2R + mg$.

3. Pour calculer R_N , on reporte l'expression de $\dot{\theta}^2$ dans l'équation projetée sur $\vec{u}_r$:

$$R_N = mg(3 \cos \theta - 2) - 2\overline{F}\theta - m \frac{v_0^2}{R}$$

4. a) Le véhicule décolle lorsque $R_N = 0$, soit pour θ_d tel que :

$$mg(3 \cos \theta_d - 2) - 2\overline{F}\theta_d - m \frac{v_0^2}{R} = 0$$

- b) Si le conducteur coupe le moteur en A, alors $\overline{F} = 0$ et θ_d vaut :

$$\theta_d = \text{Arccos} \left(\frac{2}{3} + \frac{v_0^2}{3R} \right) = 0,18 \text{ rad} = 10,85^\circ$$

Ne pas oublier de convertir la vitesse en m.s^{-1} pour l'application numérique. L'angle θ_d est inférieur à α , le véhicule décolle donc avant la fin de la descente : il ne faut donc surtout pas couper le moteur dans une descente.

- c) Pour un θ_d donné, l'expression de $\overline{F}$ est :

$$\overline{F} = \frac{1}{2\theta_d} \left(mg(3 \cos \theta_d - 2) - m \frac{v_0^2}{R} \right)$$

Pour arriver au bout de la descente, il suffit que $\theta_d = \alpha$, soit $\overline{F} = -909 \text{ N}$. Il faut donc freiner, ce à quoi on s'attendait.

Attention à convertir α en radians pour l'application numérique.

B.7 1. Le référentiel lié au sol est supposé galiléen. Les forces s'appliquant au skieur sont :

- Le poids $m \vec{g}$;
- la réaction tangentielle $\vec{T}$;
- la réaction normale du support $\vec{N}$;

- la force de frottement $\vec{F}$.

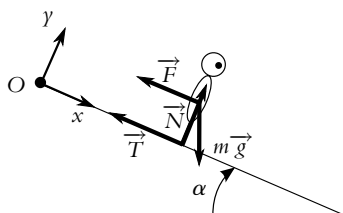


Figure S5.6

La relation fondamentale de la dynamique donne : $m\vec{a} = m\vec{g} + \vec{F} + \vec{T} + \vec{N}$, soit en projection :

$$m\ddot{x} = -T - \lambda\dot{x} + mg \sin \alpha \quad \text{et} \quad m\ddot{y} = 0 = N - mg \cos \alpha$$

On en déduit $N = mg \cos \alpha$ et donc $T = fN = fmg \cos \alpha$.

2. On reporte l'expression de T dans l'autre équation en posant $\tau = m/\lambda$ et on intègre :

$$\ddot{x} + \frac{\dot{x}}{\tau} = g(\sin \alpha - f \cos \alpha) \Rightarrow \dot{x} = A \exp\left(-\frac{t}{\tau}\right) + \tau g(\sin \alpha - f \cos \alpha)$$

Sachant qu'à $t = 0$, la vitesse est nulle, on obtient :

$$\dot{x} = \tau g(\sin \alpha - f \cos \alpha) \left(1 - \exp\left(-\frac{t}{\tau}\right)\right)$$

3. a) Lorsque $t \rightarrow \infty$, $\dot{x} \rightarrow \tau g(\sin \alpha - f \cos \alpha) = v_l$ qui constitue donc la vitesse limite. On obtient alors $\vec{v} = \vec{v}_l \left(1 - \exp\left(-\frac{t}{\tau}\right)\right)$ avec $\vec{v}_l = v_l \vec{u}_x$.
b) L'application numérique donne $v_l = 56 \text{ m.s}^{-1}$.

4. On résout l'équation $v = v_l/2$:

$$\frac{v_l}{2} = v_l \left(1 - \exp\left(-\frac{t_1}{\tau}\right)\right) \Rightarrow t_1 = 2\tau \ln 2$$

soit $t_1 = 111 \text{ s}$

5. Lorsque le skieur tombe, sa vitesse est donc $v_l/2$. L'équation du mouvement sur Oy est de la même forme que précédemment en multipliant f par 10, ce qui entraîne $T = 10fmg$. L'équation projetée sur Ox vient alors :

$$\ddot{x} = g(\sin \alpha - 10f \cos \alpha) \Rightarrow \dot{x} = g(\sin \alpha - 10f \cos \alpha)t + v_l/2$$

Le temps d'arrêt correspond à $v = 0$ soit $t_a = -v_l/(2g(\sin \alpha - 10f \cos \alpha))$. La distance d'arrêt depuis le point de chute, s'obtient en intégrant la vitesse soit :

$$d = \frac{g}{2}(\sin \alpha - 10f \cos \alpha)t_a^2 + \frac{v_l t_a}{2} \Rightarrow d = -\frac{v_l^2}{8g(\sin \alpha - 10f \cos \alpha)} = 6,2 \text{ m}$$

- B.8** 1. a) On considère la piste liée à un référentiel galiléen. Par hypothèse, les forces s'exerçant sur l'avion sont : le poids, la réaction normale de la piste et la force de traînée T due au parachute. La seule force horizontale est donc T . On appelle Ox l'axe horizontal selon lequel se déplace l'avion. La projection de la relation fondamentale de la dynamique sur Ox donne donc :

$$m\ddot{x} = -T = -k\dot{x}^2$$

avec $k = C_x \rho S / (2m)$

- b) Pour intégrer l'équation du mouvement on sépare les variables, soit en notant $v = \dot{x}$:

$$\frac{\dot{v}}{v^2} = -k \Rightarrow -\frac{1}{v} = -kt + \text{cte}$$

or à $t = 0$, $v = v_a$:

$$\frac{1}{v} = kt + \frac{1}{v_a} \Rightarrow v = \frac{v_a}{v_a kt + 1}$$

c) L'avion s'arrête si la vitesse devient nulle, or d'après l'expression précédente la vitesse tend vers 0 au bout d'un temps infini. Cette force n'est pas suffisante pour stopper l'avion.

2. On obtient l'expression de la position de l'avion en intégrant l'expression précédente de la vitesse et en prenant $x(0) = 0$:

$$x = \frac{1}{k} \ln(v_a kt + 1) + ct \text{ soit } x = \frac{1}{k} \ln(v_a kt + 1) \quad v_a kt + 1 = \exp(kx)$$

d'où la vitesse à la distance d : $v = v_a \exp(-kd) = 24,9 \text{ m.s}^{-1}$ soit une vitesse de $89,6 \text{ km.h}^{-1}$, d'où l'utilité des freins.

B.9 On considère le référentiel terrestre galiléen. Les systèmes considérés seront tour à tour le point M_1 de masse m_1 et le point M_2 de masse m_2 .

1. Les forces qui s'exercent sur M_1 et M_2 sont représentées sur les figures (S5.7 et S5.8). Les forces $\vec{R}_{N1}$ et $\vec{R}_{N2}$ représentent les résultantes des réactions normales (contacts poulie-support, poussée d'Archimède exercée par l'eau).

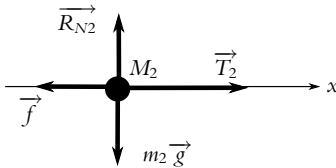


Figure S5.7

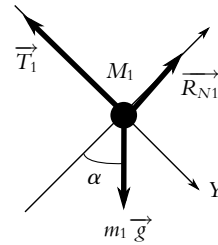


Figure S5.8

- a) La relation fondamentale de la dynamique appliquée à M_2 , en projection sur l'axe Ax donne :

$$m_2 \ddot{x} = -k\dot{x} + T_2 \quad (1)$$

- b) La relation fondamentale de la dynamique appliquée à M_1 , en projection sur l'axe CY donne :

$$m_1 \ddot{Y} = -T_1 + m_1 g \sin \alpha \quad (2)$$

- c) Les points de contact des câbles avec la poulie ont la même vitesse angulaire ω . D'autre part comme les câbles sont toujours tendus, chacun de leur point a la même vitesse que les troncs. On en déduit $\dot{x} = r_2 \omega$ et $\dot{Y} = r_1 \omega$, soit en dérivant :

$$\frac{\ddot{x}}{r_2} = \frac{\ddot{Y}}{r_1} \quad (3)$$

- d) De la relation entre les tensions, on tire $T_2 = \frac{-\Gamma + T_1 r_1}{r_2}$. On exprime T_1 avec la relation (2), puis on remplace $\ddot{Y}$ par son expression tirée de la relation (3). On reporte ces expressions dans la relation (1) et on trouve :

$$M\ddot{x} + k\dot{x} = -\frac{\Gamma}{r_2} + \frac{r_1}{r_2} m_1 g \sin \alpha$$

- e) Sans résoudre l'équation, on peut trouver la vitesse limite qui correspond à une accélération nulle (les forces motrices équilibrent alors les forces résistantes). Cela conduit à :

$$v_L = \frac{-\Gamma + r_1 m_1 g \sin \alpha}{kr_2}$$

En ce qui concerne la constante de temps on l'obtient en divisant l'équation différentielle par M ce qui conduit à une équation de la forme :

$$\ddot{x} + \frac{\dot{x}}{\tau} = \frac{-\Gamma}{Mr_2} + \frac{r_1}{Mr_2} m_1 g \sin \alpha$$

avec $\tau = M/k$. L'application numérique donne $\tau = 37,5$ s et pour v_L , il faut d'abord calculer $\alpha = \arctan(h/L_3) = 45^\circ$ soit $v_L = 0,65 \text{ m.s}^{-1}$.

- f) On peut mettre l'équation différentielle sous la forme :

$$\ddot{x} + \frac{\dot{x}}{\tau} = \frac{v_L}{\tau}$$

La solution est $\dot{x} = A \exp\left(-\frac{t}{\tau}\right) + v_L$. Au départ, les troncs d'arbres sont immobiles donc $\dot{x}(0) = 0$, ce qui entraîne : $\dot{x} = v_L \left(1 - \exp\left(-\frac{t}{\tau}\right)\right)$. On calcule le temps pour lequel $\dot{x} = 0,99v_L$, ce qui donne 173 s ou 2,9 mn.

- g) Le temps mis pour traverser le lac est : $t_2 = L_1/v_L = 21,5$ mn. On peut donc considérer que la traversée se fait à la vitesse constante v_L (sauf les trois premières minutes) et assimiler le temps de traversée à t_2 .
2. a) Les forces s'exerçant sur m_2 sont : le poids, $\vec{T}_2$ et l'on remplace $\vec{R}_{N2}$ par $\vec{R}_n$ et $\vec{f}$ par $\vec{R}_f$. La relation fondamentale de la dynamique $m_2 \vec{a} = \vec{T}_2 + \vec{R}_n + \vec{R}_f + m_2 \vec{g}$ donne en projection sur la verticale : $0 = R_n - m_2 g$, d'où l'équation projetée sur Ax :

$$m_2 \ddot{x} = T_2 - f m_2 g$$

L'équation différentielle pour Y est la même que précédemment, en revanche la relation pour la poulie donne $T_2 r_2 = T_1 r_1$. Enfin, la relation entre $\ddot{x}$ et $\ddot{Y}$ est inchangée. On obtient alors :

$$M \ddot{x} = m_1 \frac{r_1}{r_2} g \sin \alpha - m_2 f g$$

- b) L'accélération $\ddot{x}$ est donc constante : on a affaire à un mouvement uniformément accéléré. Si l'on change l'origine des temps en prenant le passage en B comme nouveau $t = 0$, sachant qu'en ce point, $\dot{x} = v_L$, on obtient :

$$\ddot{x} = \frac{g}{M} \left(m_1 \frac{r_1}{r_2} \sin \alpha - m_2 f \right) \Rightarrow \dot{x} = \frac{gt}{M} \left(m_1 \frac{r_1}{r_2} \sin \alpha - m_2 f \right) + v_L$$

et :

$$x = \frac{gt^2}{2M} \left(m_1 \frac{r_1}{r_2} \sin \alpha - m_2 f \right) + v_L t$$

- c) On obtient $\ddot{Y}$ grâce à la relation avec $\ddot{x}$, soit :

$$\ddot{Y} = \frac{r_1}{r_2} \ddot{x} = \frac{r_1}{r_2} \frac{g}{M} \left(m_1 \frac{r_1}{r_2} \sin \alpha - m_2 f \right)$$

- d) La vitesse s'annule à l'instant t_n :

$$\dot{x}(t_n) = 0 \Rightarrow t_n = -\frac{M}{g} \frac{v_L}{m_1 \frac{r_1}{r_2} \sin \alpha - m_2 f}$$

Numériquement, cette grandeur est positive ce qui montre l'existence de l'arrêt. On peut alors calculer la position d'arrêt en remplaçant t par t_n dans l'expression de x :

$$d_2 = x(t_n) = \frac{M}{2g} \frac{v_L^2}{m_2 f - m_1 \frac{r_1}{r_2} \sin \alpha}$$

L'application numérique donne $d_2 = 1$ m.

- e) En raison du rapport 2 des rayons des deux poulies, le point M_1 aura parcourue une distance deux fois plus petite que M_2 soit 150,5 m.

Chapitre 6

A.1 Le travail élémentaire de la force $\vec{f}$ s'écrit : $\delta W = \vec{f} \cdot d\vec{OM}$ avec $d\vec{OM} = dx\vec{u}_x + dy\vec{u}_y$ en coordonnées cartésiennes.

1. Pour aller de O à C directement, on suit la droite d'équation $y = \frac{y_0}{x_0}x$ soit $dy = \frac{y_0}{x_0}dx$. On en déduit :

$$W_{O \rightarrow C} = \int_0^{x_0} \left(\left(\left(\frac{y_0}{x_0} \right)^2 - 1 \right) x^2 + 3x \frac{y_0}{x_0} x \frac{y_0}{x_0} \right) dx = \left(4 \left(\frac{y_0}{x_0} \right)^2 - 1 \right) \int_0^{x_0} x^2 dx$$

$$W_{O \rightarrow C} = \frac{x_0}{3} (4y_0^2 - x_0^2)$$

2. On fait la même chose en passant par A : on décompose le chemin en OA et AC soit comme $y = 0$ et $dy = 0$ sur OA et $x = x_0$ et $dx = 0$ sur AC :

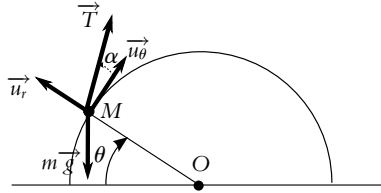
$$W_{O \rightarrow A \rightarrow C} = \int_0^{x_0} -x^2 dx + \int_0^{y_0} 3x_0 y dy = -\frac{x_0^3}{3} + 3x_0 \frac{y_0^2}{2} = \frac{x_0}{6} (9y_0^2 - 2x_0^2)$$

3. On fait la même chose en passant par B : on décompose le chemin en OB et BC soit comme $x = 0$ et $dx = 0$ sur OB et $y = y_0$ et $dy = 0$ sur BC :

$$W_{O \rightarrow B \rightarrow C} = \int_0^{y_0} 0 dy + \int_0^{x_0} (y_0^2 - x^2) dx = y_0^2 x_0 - \frac{x_0^3}{3} = \frac{x_0}{3} (3y_0^2 - x_0^2)$$

4. $W_{O \rightarrow C} \neq W_{O \rightarrow A \rightarrow C} \neq W_{O \rightarrow B \rightarrow C}$ donc le travail de $\vec{f}$ dépend du chemin suivi : il s'agit d'une force non conservative.

A.2 1. On considère le référentiel lié au sol galiléen. Le système est le personnage M . On choisit pour l'étude, une base polaire.



La force $\vec{T}$ fait avec l'horizontale (vecteur $\vec{u}_r$) un angle α . Pour un déplacement $d\vec{\ell}$ le long de la trajectoire, donc suivant $\vec{u}_r$, le travail élémentaire de la force s'écrit :

$$\delta W = \vec{T} \cdot d\vec{\ell} = T \cos \alpha d\ell$$

or en coordonnées polaires, le déplacement $d\ell$ s'écrit $d\ell = R d\theta$. On en déduit l'expression du travail sur toute la trajectoire :

$$W = \int_0^{\frac{\pi}{2}} T \cos \alpha R d\theta = \frac{\pi}{2} RT \cos \alpha$$

2. a) Les forces s'appliquant à M sont :

- le poids ;
- la réaction du support $\vec{R}_n$ (colinéaire à $\vec{u}_r$ et de même sens) ;
- le frottement du sol $\vec{F}_f$ (colinéaire à $\vec{u}_\theta$ et de sens opposé) ;
- la force de traction $\vec{T}$

La relation fondamentale donne : $m\vec{a} = m\vec{g} + \vec{T} + \vec{R}_n + \vec{F}_f$, soit en projection sur $\vec{u}_r$:

$$-mR\dot{\theta}^2 = -mg \sin \theta + R_n + T \sin \alpha \quad \text{avec} \quad v = R\dot{\theta}$$

d'où :

$$R_n = -m \frac{v^2}{R} + mg \sin \theta - T \sin \alpha$$

b) Le travail élémentaire de la force de frottement est : $\delta W = \vec{F}_f \cdot d\vec{\ell} = -\|\vec{F}_f\| R d\theta$ soit avec la relation $\|\vec{F}_f\| = f\|\vec{R}_n\|$:

$$\begin{aligned} W &= \int_0^{\frac{\pi}{2}} -f \left(-m \frac{v^2}{R} + mg \sin \theta - T \sin \alpha \right) R d\theta \\ &= Rf \left(+m \frac{\pi}{2} \frac{v^2}{R} + \frac{\pi}{2} T \sin \alpha - mg \right) \end{aligned}$$

A.3 Le référentiel lié au sol est considéré galiléen. Le système est le corps du Marsupilami auquel s'applique trois forces lorsqu'il est sur le sol :

- le poids (force conservative) ;
- la tension de la queue-ressort (force conservative) ;
- la réaction du sol (force de travail nul).

Lorsqu'il est en l'air et que la queue ne touche plus le sol, seul le poids est à prendre en compte. On en déduit que l'énergie mécanique se conserve au cours du mouvement. L'expression de l'énergie mécanique est :

$$E_m = \frac{1}{2}mv^2 + mgz + \frac{1}{2}k(z - \ell_0)^2$$

pendant que la queue touche le sol et :

$$E_m = \frac{1}{2}mv^2 + mgz$$

lorsque l'animal est en l'air ; sachant que la longueur de la queue correspond à l'altitude z de l'animal. On écrit l'égalité de l'énergie mécanique en trois points :

1. Départ : $z = \ell_m, v = 0$,
2. Décollage : $z = \ell_0, v_D$,
3. Point le plus haut : $z = h, v = 0$.

Soit :

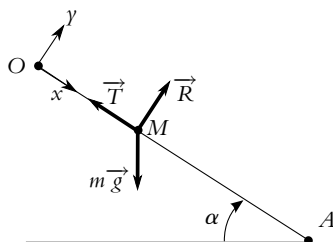
$$mg\ell_m + \frac{1}{2}k(\ell_m - \ell_0)^2 = \frac{1}{2}mv_D^2 + mg\ell_0 = mgh$$

On en déduit :

$$k = \frac{2mg(h - \ell_m)}{(\ell_m - \ell_0)^2} = 4,1 \cdot 10^3 \text{ N.m}^{-1}$$

ce qui est assez important. La vitesse est $v_D = \sqrt{2g(h - \ell_0)}$ soit $12,5 \text{ m.s}^{-1}$, ce qui est aussi une valeur importante.

A.4 1. Le référentiel terrestre est supposé galiléen. Le système M subit trois forces : le poids, la force de frottement et la réaction normale. Le poids est une force conservative, la réaction normale ne travaille pas et la force de frottement est une force non conservative.



La variation d'énergie mécanique étant égale au travail des forces non conservatives, on obtient entre le point de départ O et le point d'arrivée A :

$$E_m(A) - E_m(O) = W_{O \rightarrow A}(\vec{T})$$

Il faut donc calculer la force $\vec{T}$. La relation fondamentale de la dynamique en projection sur Oy donne $m\ddot{y} = R - mg \sin \alpha$, or $\ddot{y} = 0$ donc $R = mg \sin \alpha$.

On en déduit $\vec{T} = -fmg \sin \alpha \vec{u}_x$. Ainsi :

$$E_m(A) - E_m(O) = \int_0^{x_A} \vec{T} \cdot (dx \vec{u}_x) = -fmgL \sin \alpha$$

sachant que L (distance entre O et A) vaut $h / \cos \alpha$. On trouve finalement :

$$E_m(A) - E_m(O) = -fmgL \tan \alpha = -fmg h$$

2. En explicitant l'expression de l'énergie mécanique et en notant Z l'altitude des points, on obtient :

$$\left(\frac{1}{2} m v_A^2 + mg Z_A \right) - \left(\frac{1}{2} m v_O^2 + mg Z_O \right) = -fmg h \Rightarrow v_A = \sqrt{2gh(1-f)}$$

Application numérique : $v_A = 7,7 \text{ m.s}^{-1}$. Ne pas tenir compte du frottement revient à prendre $f = 0$ dans les expressions précédentes. On aurait alors une vitesse d'arrivée égale à $9,9 \text{ m.s}^{-1}$.

A.5 Le point A étant fixe, on le choisit comme origine du repère et on note $\vec{r}$ le vecteur $\vec{AB}$, soit aussi $r = AB$. Le vecteur unitaire est noté $\vec{u}_r$. La force $\vec{f}$ s'écrit alors $\vec{f} = \frac{q_A q_B}{4\pi\epsilon_0} \frac{\vec{u}_r}{r^2}$.

1. Pour déterminer l'énergie potentielle, on utilise la relation entre le travail et la variation d'énergie potentielle : $\delta W = -dE_p$, soit :

$$\delta W = \vec{f} \cdot d\vec{r} = \frac{q_A q_B}{4\pi\epsilon_0} \frac{\vec{u}_r}{r^2} \cdot (dr \vec{u}_r + r d\vec{u}_r).$$

Or $d\vec{u}_r$ est perpendiculaire à $\vec{u}_r$ car c'est un vecteur unitaire. On peut ainsi déterminer E_p :

$$\delta W = -dE_p = \frac{q_A q_B}{4\pi\epsilon_0 r^2} dr \Rightarrow E_p = \frac{q_A q_B}{4\pi\epsilon_0 r} + cte$$

On prend la constante nulle pour avoir une énergie potentielle nulle en l'infini.

2. Puisque $\vec{f}$ est la seule force à prendre en compte et qu'elle est conservative, l'énergie mécanique se conserve. En considérant A immobile et $q_A = q_B = q$ on peut écrire :

$$E_m = E_c + E_p = \frac{1}{2} m v_B^2 + \frac{q^2}{4\pi\epsilon_0 r}$$

Le point B s'approche de A , comme la force est répulsive, sa vitesse diminue puis s'annule, puis B s'éloigne de A . La distance minimale entre les deux points est obtenue pour $v_B = 0$. On écrit l'égalité de l'énergie mécanique entre le point de départ ($AB = r = a$) et le point où $v_B = 0$:

$$\frac{1}{2} m v_0^2 + \frac{q^2}{4\pi\epsilon_0 a} = \frac{q^2}{4\pi\epsilon_0 d_{\min}} \Rightarrow d_{\min} = \left(\frac{1}{a} + \frac{2\pi\epsilon_0 m v_0^2}{q^2} \right)^{-1}$$

3. Avec $q_A = q_B = q$, l'énergie mécanique s'écrit :

$$E_m = E_c + E_p = \frac{1}{2} m v_B^2 - \frac{q^2}{4\pi\epsilon_0 r}$$

La force est maintenant attractive. Il faut lancer B à l'opposé de A pour qu'elle puisse s'éloigner. La vitesse de B diminue, et la charge rebrousse chemin lorsque sa vitesse s'annulera. Pour éviter cela, il faut que la vitesse s'annule à l'infini, la force étant alors nulle elle-aussi. La conservation de l'énergie

mécanique donne alors :

$$\frac{1}{2}mv_0^2 - \frac{q^2}{4\pi\epsilon_0 a} = E_m(r \rightarrow \infty) = 0 \Rightarrow v_0 = \sqrt{\frac{q^2}{2\pi\epsilon_0 am}}$$

Cette vitesse porte le nom de vitesse de libération car la charge B échappe alors à l'attraction de la charge A .

B.1 Le sol est lié au référentiel galiléen terrestre. Le système « vélo+cycliste » est soumis :

- au poids (vertical) ;
- à la réaction du sol $\vec{R}$ (verticale) ;
- à la force de frottement de l'air (horizontale) ;
- à une force motrice de puissance P (horizontale).

1. Le théorème de la puissance cinétique donne :

$$\frac{dE_c}{dt} = \mathcal{P}(m\vec{g}) + \mathcal{P}(\vec{F}_f) + \mathcal{P}(\vec{R}) + P$$

Or le déplacement est horizontal donc le poids et la réaction ont une puissance nulle. La puissance de la force de frottement est :

$$\mathcal{P}(\vec{F}_f) = -kv\vec{v} \cdot \vec{v} = -kv^3$$

donc en dérivant l'énergie cinétique on trouve :

$$mv \frac{dv}{dt} = P - kv^3$$

D'autre part, on peut écrire :

$$\frac{dv}{dt} = \frac{dv}{dx} \frac{dx}{dt} = \frac{dv}{dx} v$$

En remplaçant dans la précédente équation, on trouve l'équation demandée.

2. a) La dérivée de $f(x)$ par rapport à x donne $f'(x) = -3kv^2 \frac{dv}{dx}$. L'équation différentielle précédente devient alors :

$$-\frac{m}{3k} f'(x) = f(x) \text{ ou } f'(x) + \frac{3k}{m} f(x) = 0$$

- b) On pose $\tau = \frac{m}{3k}$, ce qui donne alors : $f(x) = A \exp(-x/\tau)$. Or à $t = 0$, $x = 0$ et $v = v_0$ soit $f(0) = A = P - kv_0^3$. On en déduit finalement la vitesse :

$$v(x) = \left(\frac{1}{k} \left(P - (P - kv_0^3) \exp\left(-\frac{x}{\tau}\right) \right) \right)^{1/3}$$

B.2 1. Le référentiel lié au sol est galiléen. Les forces s'appliquant au système M sont le poids (force conservative) et la réaction du support $\vec{R}$ (force de travail nul).

Pour calculer la vitesse en O , on peut soit appliquer le théorème de l'énergie cinétique, soit appliquer le théorème de l'énergie mécanique.

On résout en premier lieu avec le théorème de l'énergie cinétique :

$$E_c(O) - E_c(A) = W_{A \rightarrow O}(m\vec{g}) + W_{A \rightarrow O}(\vec{R}) = W_{A \rightarrow O}(m\vec{g})$$

soit puisque $v_A = 0$:

$$\frac{1}{2}mv_O^2 = mg(y_A - y_O) = mgh \Rightarrow v_O = \sqrt{2gh}$$

Si on résout avec le théorème de l'énergie mécanique, on utilise le fait que le poids est conservatif et que la réaction normale ne travaille pas, donc l'énergie mécanique se conserve :

$$E_m(A) = E_m(O) \Rightarrow \frac{1}{2}mv_A^2 + mgy_A = \frac{1}{2}mv_O^2 + mgy_O$$

On retrouve le résultat précédent.

Pour calculer la vitesse en un point M quelconque du cercle, on applique l'un des deux théorèmes précédents soit entre A et M , soit entre O et M . D'après la figure (S6.1) on a l'altitude de M : $z = a(1 - \cos \theta)$. On obtient alors :

$$v_m = \sqrt{2(h + a(\cos \theta - 1))}$$

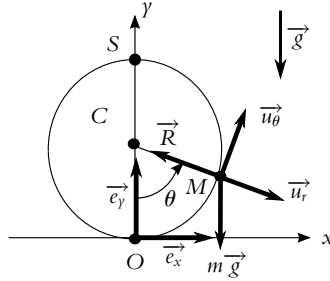


Figure S6.1

2. Pour déterminer la réaction, on applique la relation fondamentale de la dynamique en projection sur $\vec{u}_r$ soit :

$$-ma\dot{\theta}^2 = -R + mg \cos \theta$$

or $\dot{\theta} = v_M/a$. En utilisant la vitesse v_M calculée précédemment, on obtient :

$$R = mg \left(3 \cos \theta + 2 \frac{h}{a} - 2 \right)$$

3. Pour que la bille ait un mouvement révolutif, il faut que la réaction ne s'annule en aucun point du cercle. Son expression montre qu'elle est minimale en $\theta = \pi$ (au sommet S du cercle). On souhaite donc que $R(\theta = \pi) > 0$ soit :

$$mg \left(-3 + 2 \frac{h}{a} - 2 \right) > 0 \Rightarrow h > \frac{5}{2}a$$

Ce résultat est contraire à l'intuition qui nous laisserait penser qu'il suffit que la bille soit lâchée d'une hauteur légèrement supérieure à celle de S , mais cette condition n'assurerait qu'une vitesse non nulle en S .

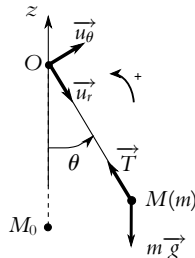
- B.3** 1. Le référentiel du laboratoire est considéré comme galiléen. Le pendule est soumis à deux forces : son poids (conservatif) et la tension du fil qui ne travaille pas. L'énergie mécanique se conserve donc. On prend comme origine de l'énergie potentielle, le point d'attache O du pendule soit :

$$E_m = \frac{1}{2}mv^2 + mgz$$

Puisque l'énergie mécanique se conserve, on écrit l'égalité de cette énergie en $M(z, v)$ et le point le plus bas $M_0(-\ell, v_0)$, soit :

$$\frac{1}{2}mv^2 + mgz = \frac{1}{2}mv_0^2 - mg\ell \Rightarrow \|\vec{v}\| = \sqrt{v_0^2 - 2g(z + \ell)}$$

- 2.



La relation fondamentale de la dynamique en projection sur $\vec{u}_r$ donne :

$$-T + mg \cos \theta = -m\ell \dot{\theta}^2$$

Or $v = \ell \dot{\theta}$ et $z = -\ell \cos \theta$ d'où :

$$T = m \frac{v^2}{\ell} - mg \frac{z}{\ell}$$

En remplaçant v^2 par son expression, on obtient :

$$T = \frac{m}{\ell} (v_0^2 - 3gz - 2g\ell)$$

3. La première étude à faire est de déterminer les altitudes d'annulation de T et v . La tension T s'annule pour $z_T = \frac{v_0^2 - 2g\ell}{3g}$. La vitesse v s'annule pour $z_v = \frac{v_0^2 - 2g\ell}{2g}$.

Première conclusion : si $v_0^2 > 2g\ell$, la tension s'annule avant la vitesse ($0 < z_T < z_v$) et si $v_0^2 < 2g\ell$ c'est le contraire. Ainsi si $v_0^2 < 2g\ell$, le mouvement est oscillatoire et le fil reste tendu.

Deuxième étude : il faut savoir si la vitesse ou la tension s'annulent avant que M ne passe au point supérieur ($\theta = \pi, z = \ell$). Les conditions de non annulation de T et v en ce point sont :

- pour v , il faut $z_v > \ell$ soit $v_0^2 > 4g\ell$;
- pour T , il faut $z_T > \ell$ soit $v_0^2 > 5g\ell$.

Ainsi si $v_0^2 > 5g\ell$, ni la tension ni la vitesse ne s'annulent. Le pendule fait donc un tour complet sans se détendre : le mouvement est révolutif ininterrompu.

En revanche, si $v_0^2 < 5g\ell$, la tension s'annule avant que M n'atteigne le point le plus haut, le fil se détend et le mouvement est donc interrompu. À priori, lorsque $v_0^2 < 4g\ell$, le pendule n'a pas assez d'énergie pour faire un tour et le mouvement est interrompu oscillatoire ; pour être sûr il faudrait étudier le mouvement alors uniquement dans le champ de pesanteur. Si $v_0^2 > 4g\ell$ le mouvement est interrompu révolutif.

Récapitulatif :

- $v_0^2 < 2g\ell$: mouvement oscillatoire ininterrompu ;
- $2g\ell < v_0^2 < 4g\ell$: mouvement oscillatoire interrompu ;
- $4g\ell < v_0^2 < 5g\ell$: mouvement révolutif interrompu ;
- $v_0^2 > 5g\ell$: mouvement révolutif non interrompu.

B.4 Le référentiel du laboratoire galiléen est repéré par un axe Ox . Le système est la charge q qui se déplace uniquement selon l'axe Ox .

1. a) Pour déterminer l'énergie potentielle on calcule le travail élémentaire :

$$\delta W = \vec{f} \cdot d\vec{\ell} = \frac{q^2}{4\pi\epsilon_0} \frac{dx}{x^2} = -dE_p \Rightarrow E_p = \frac{q^2}{4\pi\epsilon_0} \frac{1}{x} + cte$$

d'après l'énoncé, on choisit la constante nulle.

- b) La seule force exercée sur M dérive d'une énergie potentielle, elle est donc conservative. On écrit l'égalité de l'énergie mécanique entre la position initiale ($x = a, v_0 = 0$) et la position finale ($x \rightarrow +\infty, v_f$) :

$$0 + \frac{q^2}{4\pi\epsilon_0} \frac{1}{a} = \frac{1}{2}mv_f^2 + 0 \Rightarrow v_f = \sqrt{\frac{q^2}{2\pi\epsilon_0 m} \frac{1}{a}}$$

2. La force résultante est la somme des forces exercées par chaque charge en O et O' sur M , soit :

$$\vec{f} = \left(\frac{q^2}{4\pi\epsilon_0} \frac{1}{OM^2} - \frac{4q^2}{4\pi\epsilon_0} \frac{1}{O'M^2} \right) \vec{u}_x$$

or $OM = x$ et $O'M = 2a - x$, soit :

$$\vec{f} = \frac{q^2}{4\pi\epsilon_0} \left(\frac{1}{x^2} - \frac{4}{(2a-x)^2} \right) \vec{u}_x$$

La position d'équilibre est obtenue pour $\vec{f} = \vec{0}$, donc :

$$(2a - x)^2 = 4x^2 \Rightarrow 3x^2 + 4ax - 4a^2 \Rightarrow x = -2a \text{ ou } x = \frac{2}{3}a$$

La solution physique est celle appartenant à $[0, 2a]$, donc la position d'équilibre est $x_e = \frac{2}{3}a$.

3. L'énergie potentielle totale est la somme des énergies potentielles des deux forces soit :

$$E_p = \frac{q^2}{4\pi\epsilon_0} \left(\frac{1}{x} + \frac{4}{2a - x} \right)$$

Attention dans le deuxième terme, c'est la distance entre les charge en O' et M donc il n'y'a pas le signe « - » de la force. Pour s'en convaincre, on peut dériver cette énergie potentielle et on retrouvera la bonne expression de la force. On dérive deux fois l'énergie potentielle : le signe de la dérivée seconde indique la stabilité :

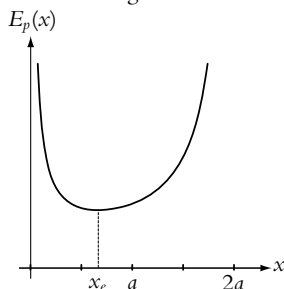
$$\frac{dE_p}{dx} = \frac{q^2}{4\pi\epsilon_0} \left(-\frac{1}{x^2} + \frac{4}{(2a - x)^2} \right) \text{ et } \frac{d^2E_p}{dx^2} = \frac{q^2}{4\pi\epsilon_0} \left(\frac{2}{x^3} - \frac{8}{(2a - x)^3} \right)$$

La valeur de la dérivée seconde pour x_e est $\frac{q^2}{2\pi\epsilon_0} \frac{27}{16a^3}$ donc positive. La position d'équilibre correspond à un minimum d'énergie potentielle, elle est donc stable.

4. L'expression de l'énergie mécanique est :

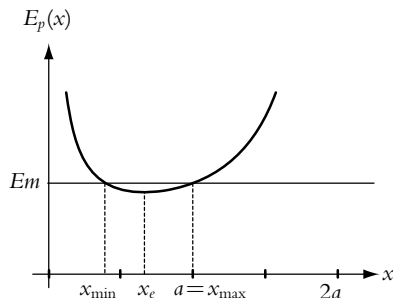
$$E_m = \frac{1}{2}m\dot{x}^2 + \frac{q^2}{4\pi\epsilon_0} \left(\frac{1}{x} + \frac{4}{2a - x} \right)$$

5. L'allure de la courbe $E_p(x)$ est donnée sur la figure :



Le minimum a pour valeur $E_p(x_e)$ soit $\frac{9q^2}{8\pi\epsilon_0 a}$.

6. La zone de mouvement possible est telle que $E_m \geq E_p$ puisque $E_m = E_c + E_p$ et que $E_c \geq 0$. Initialement, le point M est lâché en $x = a$ sans vitesse initiale, donc $E_m = E_p(x = a)$. On a tracé sur la courbe d'énergie potentielle la droite $E_p = E_m$. Les zones accessibles correspondent donc aux points situés sous cette droite. Les intersections avec la courbe $E_p(x)$ correspondent aux extrema de x pour lesquels $v = \dot{x} = 0$. On en déduit que $x_{\max} = a$ et $0 < x_{\min} < a/2$.



L'énergie cinétique étant la différence entre l'énergie mécanique et l'énergie potentielle, elle est maximale lorsque l'énergie potentielle est minimale donc pour $x = x_e$, soit :

$$E_{c,max} = E_m - E_p(x_e) = E_p(a) - E_p(x_e) = \frac{q^2}{4\pi\epsilon_0} \left(\frac{1}{a} + \frac{4}{2a-a} - \frac{9}{2a} \right) = \frac{q^2}{8\pi\epsilon_0 a}$$

Un calcul permet de trouver $x_{\min}$ et $x_{\max}$. On cherche par le calcul les positions telles que $E_m = E_p(x)$, soit :

$$E_p(x) = E_p(a) \Rightarrow \frac{1}{x} + \frac{4}{2a-x} = \frac{5}{a} \Rightarrow 5x^2 - 7ax + 2a^2 = 0 \Rightarrow x_{\max} = a \text{ et } x_{\min} = \frac{2}{5}a$$

Pour passer des fractions à l'équation du second degré, on a réduit tous les termes au même dénominateur.

7. La relation fondamentale de la dynamique donne $m\vec{a} = \vec{f}$ avec $m\vec{a} = m\ddot{x}\vec{u}_x$ et $\vec{f}$ donnée précédemment. Le mouvement a lieu au voisinage de la position d'équilibre donc $x = x_e + \epsilon$ avec $\epsilon \ll x_e$. On a alors $\ddot{x} = \ddot{\epsilon}$. Pour la force, on effectue un développement limité :

$$f(x) = \frac{q^2}{4\pi\epsilon_0} \left(\frac{1}{x^2} - \frac{4}{(2a-x)^2} \right) = \frac{q^2}{4\pi\epsilon_0} \left(\frac{1}{(x_e + \epsilon)^2} - \frac{4}{(2a - x_e - \epsilon)^2} \right)$$

On met x_e en facteur dans le premier terme et $2a - x_e$ dans le second :

$$f(\epsilon) = \frac{q^2}{4\pi\epsilon_0} \left(\frac{1}{x_e^2} \left(1 + \frac{\epsilon}{x_e} \right)^{-2} - \frac{4}{(2a - x_e)^2} \left(1 + \frac{\epsilon}{(2a - x_e)} \right)^{-2} \right)$$

En utilisant la formule du développement au premier ordre :

$$f(\epsilon) = \frac{q^2}{4\pi\epsilon_0} \left(\frac{1}{x_e^2} \left(1 - 2\frac{\epsilon}{x_e} \right) - \frac{4}{(2a - x_e)^2} \left(1 - 2\frac{\epsilon}{(2a - x_e)} \right) \right)$$

En remplaçant x_e par sa valeur on trouve finalement $f(\epsilon) = -\frac{q^2}{4\pi\epsilon_0} \frac{81}{4a^3} \epsilon$. On obtient alors l'équation différentielle des petites oscillations :

$$m\ddot{\epsilon} = f(\epsilon) \Rightarrow \ddot{\epsilon} + \omega_0^2 \epsilon = 0 \text{ avec } \omega_0 = \sqrt{\frac{q^2}{4\pi\epsilon_0} \frac{81}{4ma^3}}$$

- B.5** 1. Le référentiel du laboratoire est considéré comme galiléen. Le système M est soumis à son poids, à la tension du ressort et à la réaction de la tige.

La réaction ne travaille pas et les deux autres forces sont conservatives donc l'énergie mécanique est constante.

L'énergie potentielle associée au poids est $mgz + cte$ et celle dont dérive la tension du ressort $\frac{1}{2}k(\ell - \ell_v)^2 + cte$. Dans le cas présent, la longueur du ressort $\ell = X$ et l'altitude $z = X \cos \alpha$. D'où l'expression de l'énergie potentielle totale :

$$E_p = mgX \cos \alpha + \frac{1}{2}k(X - \ell_v)^2 + cte$$

On prend par exemple $E_p = 0$ pour $X = \ell_v$, d'où l'expression de la constante : $cte = -mg\ell_v \cos \alpha$.

2. On exprime l'énergie mécanique :

$$E_m = \frac{1}{2}m\dot{X}^2 + mgX \cos \alpha + \frac{1}{2}k(X - \ell_v)^2 - mg\ell_v \cos \alpha$$

Comme elle est constante, sa dérivée est nulle :

$$\frac{dE_m}{dt} = \dot{X}(m\ddot{X} + mg \cos \alpha + k(X - \ell_v)) = 0$$

On exclut $\dot{X} = 0 \forall t$ qui correspond à l'immobilité, d'où l'équation du mouvement :

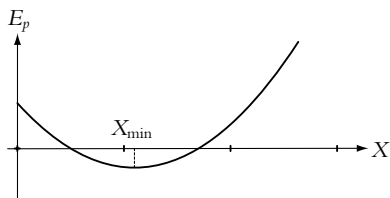
$$\ddot{X} + \frac{k}{m}X = \frac{k}{m}\ell_v - mg \cos \alpha$$

3. a) Pour étudier la fonction E_p , on commence par calculer sa dérivée :

$$\frac{dE_p}{dX} = mg \cos \alpha + k(X - \ell_v)$$

Cette dérivée est nulle pour $X = \ell_v - \frac{mg \cos \alpha}{k}$ qui est bien positif d'après l'hypothèse de l'énoncé.

Cette position correspond à un minimum $E_{p,\min} = -\frac{(mg \cos \alpha)^2}{2k}$.



- b) Comme l'énergie cinétique est positive ou nulle, le mouvement a lieu dans les zones telles que $E_m \geq E_p$, l'égalité correspondant aux positions extrêmes. Sur la figure (S6.2), on a tracé l'énergie mécanique égale à $\frac{1}{2}V_0^2$, car à $t = 0$ $E_p = 0$. Elle est supérieure à E_p en ℓ_v car l'énergie cinétique initiale n'est pas nulle.

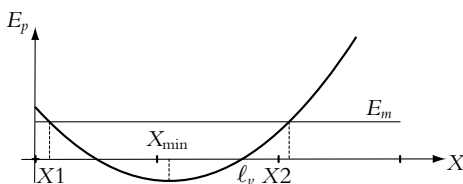


Figure S6.2

Les points extrêmes atteints correspondent à une énergie cinétique nulle, donc :

$$E_m = E_p \Rightarrow \frac{1}{2}mV_0^2 = mgX \cos \alpha + \frac{1}{2}k(X - \ell_v)^2 - mg\ell_v \cos \alpha$$

D'après la figure $X_1 > 0$ si $E_m < E_p(0)$ soit :

$$\frac{1}{2}mV_0^2 < \frac{1}{2}k\ell_v^2 - mg\ell_v \cos \alpha \Rightarrow V_0 < \pm \sqrt{\frac{k\ell_v^2 - 2mg\ell_v \cos \alpha}{m}}$$

En pratique, le point O ne pourra pas être atteint à cause de la longueur minimale du ressort lorsque toutes les spires se touchent.

- c) On obtient l'expression de la vitesse en écrivant que l'énergie mécanique à tout instant est égale sa valeur à $t = 0$:

$$\frac{1}{2}mV^2 + mgX \cos \alpha + \frac{1}{2}k(X - \ell_v)^2 - mg\ell_v \cos \alpha = \frac{1}{2}mV_0^2$$

soit

$$V = \pm \sqrt{V_0^2 - 2gX \cos \alpha - \frac{k(X - \ell_v)^2}{m} + 2g\ell_v \cos \alpha}$$

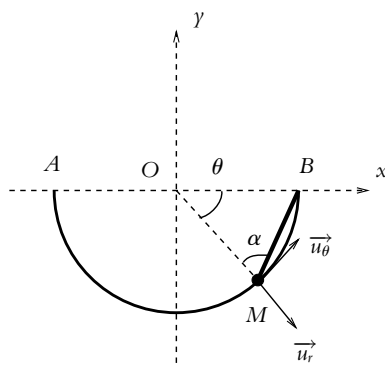
Pour déterminer $X(t)$, on utilise l'équation du mouvement établie à la question 2. On pose $X_{\min} = \ell_v - \frac{mg \cos \alpha}{k}$ et $\omega_0 = \sqrt{k/m}$, d'où :

$$\ddot{X} + \omega_0^2 X = \omega_0^2 X_{\min} \Rightarrow X = X_{\min} + A \cos \omega_0 t + B \sin \omega_0 t$$

On détermine les constantes à $t = 0$ soit, $X(0) = \ell_v = X_{\min} + A$ et $\dot{X}(0) = V_0 = \omega_0 B$ d'où :

$$X(t) = X_{\min} + (\ell_v - X_{\min}) \cos \omega_0 t + \frac{V_0}{\omega_0} \sin \omega_0 t$$

B.6 1. Le dispositif étudié est :



- a) Le triangle OMB est isocèle en O donc les angles des sommets B et M sont égaux. Par ailleurs, la somme des angles d'un triangle vaut π . En explicitant ces deux conditions, on obtient $\alpha = \frac{\pi - \theta}{2}$.

- b) Pour calculer la distance MB , on détermine son carré

$$MB^2 = (\overrightarrow{MO} + \overrightarrow{OB})^2 = R^2 + R^2 + 2R^2 \cos(\pi - \theta) = 4R^2 \sin^2 \frac{\theta}{2}$$

On en déduit $MB = 2R \left| \sin \frac{\theta}{2} \right|$.

- c) On étudie la bille dans le référentiel terrestre galiléen. Elle est soumise à son poids, à la tension $\vec{T}$ du ressort et à la réaction $\vec{N}$ du cerceau qui est normale du fait de l'absence de frottement. On projette le principe fondamental sur la tangente au cercle ou sur $\vec{u}_\theta$ en coordonnées polaires (il faut faire attention au fait que la tension de l'élastique n'est pas dirigée sur cette direction) :
- $$mR\ddot{\theta} = 2kR \sin \frac{\theta}{2} \cos \frac{\theta}{2} - mg \cos \theta \text{ soit } \ddot{\theta} = \frac{k}{m} \sin \theta - \frac{g}{R} \cos \theta.$$

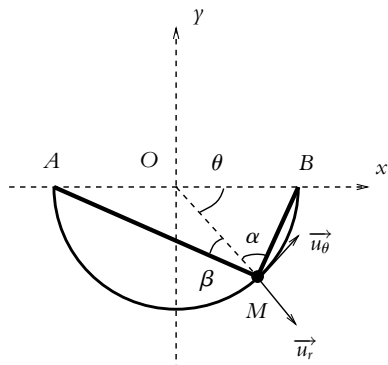
- d) Il y a équilibre lorsque $\tan \theta_e = \frac{mg}{ka}$.

- e) Au voisinage de la position d'équilibre, on a $\theta \simeq \theta_e + \epsilon$ et $\ddot{\theta} = \ddot{\epsilon}$. En effectuant un développement limité de l'équation du mouvement, on en déduit :

$$\ddot{\epsilon} - \left(\frac{g}{a} \sin \theta_e + \frac{k}{m} \cos \theta_e \right) \epsilon = \ddot{\epsilon} - \cos \theta_e \left(\frac{k}{m} + \frac{mg^2}{ka^2} \right) \epsilon = 0.$$

- f) On aura des oscillations si $\cos \theta_e > 0$ et la position d'équilibre sera stable.

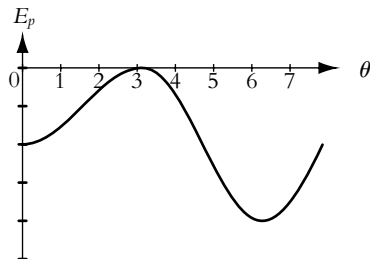
2. On considère maintenant le système :



- a) Par le même raisonnement que pour le premier système, on obtient $\beta = \frac{\theta}{2}$.
- b) De même, pour la distance MA , on a $MA = 2R \left| \cos \frac{\theta}{2} \right|$.
- c) On étudie toujours la bille dans le référentiel terrestre galiléen. Elle est soumise à son poids, aux tensions $\vec{T}_1$ et $\vec{T}_2$ des ressorts et à la réaction $\vec{N}$ du cerceau qui est normale du fait de l'absence de frottement. La projection du principe fondamental de la dynamique sur la tangente au cercle ou sur $\vec{u}_\theta$ en coordonnées polaires donne $mR\ddot{\theta} = -mg \cos \theta + 2kR \sin \frac{\theta}{2} \sin \alpha - 2kR \cos \frac{\theta}{2} \sin \beta = -mg \cos \theta$.
- d) Les ressorts n'ont pas d'influence sur θ car la résultante des tensions des ressorts est nulle sur cette direction.
- e) La résultante des tensions des ressorts s'écrit $-k\vec{AM} - k\vec{BM} = -2k\vec{OM}$ donc on a la même conclusion.
- f) En appliquant la définition de l'énergie potentielle comme l'opposé du travail $\delta W = -dE_p$, on calcule le travail élémentaire des forces et on détermine la constante avec le choix de l'origine proposée. On obtient $E_p = mgR(1 + \sin \theta)$.
- g) L'énergie mécanique $E_m = E_c + E_p$ est conservée du fait que les forces dérivent d'une énergie potentielle ou ne travaillent pas. On peut donc établir l'équation du mouvement en dérivant l'expression de l'énergie mécanique par rapport au temps. Or $E_m = mgR(1 + \sin \theta) + \frac{1}{2}mR\dot{\theta}^2$, on retrouve bien l'équation du mouvement trouvée précédemment.
- h) On a une position d'équilibre pour $\frac{dE_p}{d\theta} = 0$ soit $mgR \cos \theta = 0$. On en déduit que les positions d'équilibre correspondent à $\theta = \pm \frac{\pi}{2}$.
- i) Leur stabilité dépend du signe de $\frac{d^2E_p}{d\theta^2} = -mgR \sin \theta$. Cette quantité est positive et l'équilibre est stable pour $\theta = -\frac{\pi}{2}$. La dérivée seconde est négative et l'équilibre instable pour $\theta = \frac{\pi}{2}$.
- j) Le principe fondamental de la dynamique projeté sur la normale au cercle donne $-mR\dot{\theta}^2 = mg \sin \theta - N - 2kR$. En utilisant la conservation de l'énergie mécanique E_m , on peut exprimer $\dot{\theta}^2$ par $R\dot{\theta}^2 = 2g(\sin \theta_0 - \sin \theta)$. On en déduit l'expression de la réaction normale $N = mg(2 \sin \theta_0 - \sin \theta) - 2kR$.
- k) Il y a décollage si la réaction normale s'annule soit $N = 0$ ou $\sin \theta = 2 \sin \theta_0 - \frac{2kR}{mg}$.
- l) Comme pour toute valeur de θ , on a $-1 \leq \sin \theta \leq 1$. Le décollage est envisageable si $2kR \leq mg$.

B.7 Le référentiel du laboratoire est supposé galiléen.

- On prend un axe Oz vertical ascendant d'origine C_2 . L'énergie potentielle est alors donnée par l'expression $E_p = mgz + \text{cte}$. L'énoncé propose de prendre l'énergie potentielle nulle en B soit avec l'origine choisie, pour $z = R_2$, d'où l'expression de l'énergie $E_p = mg(z - R_2)$.
On cherche maintenant à exprimer E_p en fonction de θ . Pour la partie (1), on peut écrire $z = R_2 - R_1 - R_1 \cos \theta$ d'où $E_{p1} = -mgR_1(1 + \cos \theta)$.
Pour la partie (2), on trouve $z = -R_2 \cos \theta$, soit $E_{p2} = -mgR_2(1 + \cos \theta)$.
- Les résultats précédents permettent de tracer la courbe $E_p(\theta)$.



- Les positions angulaires d'équilibre apparaissent sur la figure précédente : ce sont les minima (positions stables) en $\theta = 0$ et 2π et le maximum en $\theta = \pi$. On les retrouve par le calcul : $\frac{dE_p}{d\theta} = mgR_i \sin \theta$. L'indice i étant égal à 1 ou 2 suivant la zone. On trouve donc des extrema pour $\sin \theta = 0$. Dans l'intervalle $\left[-\frac{\pi}{2}, \frac{5\pi}{2}\right]$, les solutions sont $\theta \in 0, \pi, 2\pi$. La dérivée seconde de E_p est égale à $mgR_i \cos \theta$, positive pour $\theta = 0$ et 2π correspondant aux positions stables et négative pour $\theta = \pi$ correspondant à la position instable.
- a) Pour atteindre le puits de potentiel en F depuis A , l'anneau doit passer la barrière de potentiel en B . Un mouvement est possible si l'énergie mécanique vérifie $E_m \geq E_p$, il faut donc donner au point M lorsqu'il est en A une énergie mécanique au moins égale à $E_p(B) = 0$ soit :

$$E_m(A) = E_c(A) + E_p(A) = \frac{1}{2}mV_0^2 - mgR_1$$

d'où

$$E_m(A) \geq 0 \Rightarrow V_0 \geq \sqrt{2gR_1}$$

- On écrit l'égalité de l'énergie mécanique en A et en F :

$$\frac{1}{2}mV_0^2 - mgR_1 = \frac{1}{2}mV_F^2 - 2mgR_2 \Rightarrow V_F = \sqrt{V_0^2 - 2g(R_1 - 2R_2)}$$

- Pour sortir de l'anneau en S , il faut que l'anneau ait une énergie mécanique supérieure à $E_p(S)$, ce qui est le cas s'il franchit la barrière de potentiel en B . Donc la condition est la même que pour atteindre F .

B.8 Le référentiel du laboratoire est considéré galiléen. Les forces s'exerçant sur le système sont : le poids $m\vec{g}$ et la tension du fil $\vec{T}$.

- Le théorème de la puissance mécanique appliqué à la première phase donne :

$$\frac{dE_c}{dt} = \mathcal{P} \Rightarrow m\vec{v} \cdot \frac{d\vec{v}}{dt} = m\vec{g} \cdot \vec{v} + \vec{T} \cdot \vec{v}$$

Sachant que pendant la première phase, le mouvement est circulaire de rayon L , on a $\vec{v} = L\dot{\theta}\vec{u}_\theta$. Ainsi puisque $\vec{T}$ est perpendiculaire à $\vec{v}$, on trouve :

$$L^2\ddot{\theta} = -mgL\dot{\theta} \sin \theta$$

En excluant l'immobilité ($\dot{\theta} = 0 \forall t$), on trouve l'équation classique du pendule :

$$\ddot{\theta} + \frac{g}{L} \sin \theta = 0$$

2. Pour $\sin \theta \simeq \theta$, on trouve des oscillations sinusoïdales de pulsation $\omega_0 = \sqrt{g/L}$ et de période $T_0 = 2\pi/\omega_0$. Entre l'instant initial et t_1^- il y a un quart de période, soit $T_0/4$, d'où $\delta t_I = \frac{\pi}{2} \sqrt{\frac{L}{g}}$.

3. La tension ne travaille pas et le poids est conservatif donc l'énergie mécanique E_m se conserve. L'énergie potentielle du poids a pour expression $mgz = -mgL \cos \theta$ si on la prend nulle en O . On écrit l'égalité des expressions de E_m à $t = 0$ et à t_1^- :

$$-mgL \cos \theta_0 = \frac{1}{2} m v_1^{-2} - mgL \rightarrow v_1^- = -\sqrt{2gL(1 - \cos \theta_0)}$$

La vitesse angulaire est donc $\omega_1^- = v_1^-/L$ soit $\omega_1^- = -\sqrt{\frac{2g(1 - \cos \theta_0)}{L}}$.

4. D'après l'énoncé $E_m(t_1^-) = E_m(t_1^+)$, donc puisque la position est pratiquement la même, l'énergie potentielle est la même et il y a égalité des énergies cinétique d'où $v_1^- = v_1^+$. En revanche $\omega_1^+ = v_1^+/(2L/3)$ d'où :

$$\omega_1^+ = -\frac{3}{2} \sqrt{\frac{2g(1 - \cos \theta_0)}{L}}$$

5. La deuxième phase correspond à une demi période d'un pendule de longueur $2L/3$, donc la durée est $\delta t_{II} = \pi \sqrt{\frac{2L}{3g}}$.

6. On écrit de nouveau l'égalité de l'énergie mécanique entre t_1^+ et t_2 sachant qu'à t_2 l'altitude z vaut $-\frac{L}{3} - \frac{2L}{3} \cos \theta_2$:

$$\frac{1}{2} m v_1^{+2} - mgL = -mg \frac{L}{3} (1 + 2 \cos \theta_2) \Rightarrow \cos \theta_2 = \frac{3 \cos \theta_0 - 1}{2}$$

7. Après la phase 2 on retrouve la phase 1 et donc le pendule atteint l'angle initial θ_0 pour lequel la vitesse s'annule puis, de nouveau décroissance de θ jusqu'à 0, et de nouveau une phase 2. Dans la réalité l'oscillation va peu à peu s'amortir en raison des frottements.

La période T est la somme d'une demi période d'un pendule de longueur L et d'une demi période

d'un pendule de longueur $2L/3$ soit $T = \pi \sqrt{\frac{L}{g}} \left(1 + \sqrt{\frac{2}{3}} \right)$.

8. Le portrait de phase d'un oscillateur sinusoïdal est une ellipse, donc ici on aura deux demi ellipses :

$$\text{phase 1} \quad \begin{cases} \theta = \theta_0 \cos \omega_0 t \\ \dot{\theta} = -\omega_0 \theta_0 \sin \omega_0 t \end{cases} \quad \text{phase 2} \quad \begin{cases} \theta = \theta_2 \cos \left(\sqrt{\frac{3}{2}} \omega_0 t + \varphi \right) \\ \dot{\theta} = -\sqrt{\frac{3}{2}} \omega_0 \theta_2 \sin \left(\sqrt{\frac{3}{2}} \omega_0 t + \varphi \right) \end{cases}$$

La discontinuité de la vitesse angulaire est due à la discontinuité de la longueur du fil.

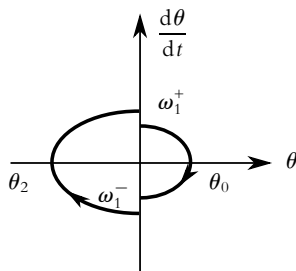


Figure S6.3

Chapitre 7

On rappelle que de manière générale la force de tension exercée par un ressort est $\vec{T} = -k(\ell - \ell_0)\vec{u}_{\text{ext}}$ où ℓ est la longueur du ressort, ℓ_0 sa longueur à vide et $\vec{u}_{\text{ext}}$ un vecteur unitaire orienté du point d'attache immobile vers la masse dont on étudie le mouvement.

A.1 1. Pour déterminer les caractéristiques du ressort équivalent aux deux ressorts en parallèle, on va écrire l'équation du mouvement de la masse m soumise aux deux tensions et à son poids (figure S7.1). On suppose que la masse reste horizontale et que les deux ressorts ont toujours la même longueur. Les grandeurs relatives au ressort de gauche (respectivement de droite) sont indicées par 1 (respectivement 2). On note $\vec{u}_z$ le vecteur unitaire de l'axe Oz . Dans le cas présent $\vec{u}_{\text{ext}} = \vec{u}_z$ et les longueurs ℓ_1 et ℓ_2 sont égales à z .

La relation fondamentale de la dynamique donne : $m\vec{a} = m\vec{g} + \vec{T}_1 + \vec{T}_2$ soit en projection :

$$m\ddot{z} = -mg - k(z - \ell_0) - k(z - \ell_0) = -mg - 2k(z - \ell_0)$$

Le ressort équivalent a donc une longueur à vide $\ell_{e0} = \ell_0$ et une constante de raideur $k_e = 2k$.

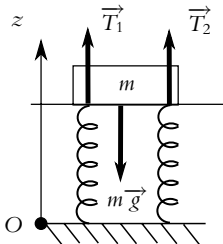


Figure S7.1

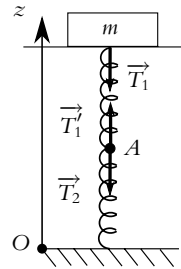


Figure S7.2

2. Les grandeurs relatives au ressort du haut (respectivement du bas) sont indicées par 1 (respectivement 2). On note $\vec{u}_z$ le vecteur unitaire de l'axe Oz . Sur la figure (S7.2) on n'a pas représenté le poids pour plus de clarté. Les forces s'exerçant sur la masse m sont : son poids $m\vec{g}$ et la tension $\vec{T}_1 = -k_1(\ell_1 - \ell_{01})\vec{u}_{\text{ext},1}$. Les forces s'exerçant sur le point A sont les deux tensions $\vec{T}_2 = -k_1(\ell_2 - \ell_{02})\vec{u}_{\text{ext},2}$ et $\vec{T}_1' = -\vec{T}_1$ (les ressorts sont de masse négligeable).

Comment définir le point fixe pour le ressort 1 ? Cela dépend du système étudié : si le système est la masse m alors c'est comme si le point fixe est A ; dans le cas où le système est A alors c'est comme si le point fixe est la masse m .

D'après ce qui précède $\vec{u}_{\text{ext},1} = \vec{u}_{\text{ext},2} = \vec{u}_z$ et $\vec{u}'_{\text{ext},1} = -\vec{u}_z$.

La relation fondamentale appliquée à m en projection sur Oz :

$$m\ddot{z} = -k_1(\ell_1 - \ell_{01}) - mg$$

et pour A (masse m_A nulle) :

$$m_A\ddot{z}_A = 0 = k_1(\ell_1 - \ell_{01}) - k_2(\ell_2 - \ell_{02}) \Rightarrow \ell_2 = \frac{k_1}{k_2}(\ell_1 - \ell_{01}) + \ell_{02}$$

D'autre part le ressort équivalent que l'on recherche doit donner les mêmes positions de m et le même mouvement. On en déduit que sa longueur est égale à $\ell_1 + \ell_2$. L'équation du mouvement de m attaché à ce ressort équivalent est donc :

$$m\ddot{z} = -k_e(\ell_e - \ell_{0e}) - mg = -k_e(\ell_1 + \ell_2 - \ell_{0e}) - mg$$

En remplaçant ℓ_2 par l'expression calculée précédemment, on obtient :

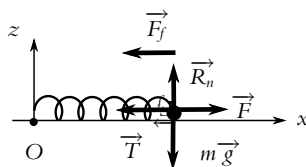
$$m\ddot{z} = -k_e \left(\left(1 + \frac{k_1}{k_2} \right) \ell_1 - \frac{k_1}{k_2} \ell_{01} + \ell_{02} - \ell_{0e} \right) - mg$$

Les deux équations du mouvement de m doivent être semblables pour tout ℓ_1 . On en déduit :

1. $k_1 = k_e \left(1 + \frac{k_1}{k_2} \right)$;
2. $k_1 \ell_{01} = k_e \left(\frac{k_1}{k_2} \ell_{01} - \ell_{02} + \ell_{0e} \right)$.

La première relation donne $k_e = \frac{k_1 k_2}{k_1 + k_2}$ et la deuxième $\ell_{0e} = \ell_{01} + \ell_{02}$.

A.2 1. Dans le référentiel galiléen du laboratoire, le système M est soumis à son poids, la réaction du support horizontal $\vec{R}_n$, la tension du ressort $\vec{T}$, la force de frottement et la force constante $\vec{F}$. On note $\vec{u}_x$ le vecteur unitaire de l'axe Ox .



La relation fondamentale de la dynamique donne : $m \vec{a} = m \vec{g} + \vec{R}_n + \vec{T} + \vec{F}_f + \vec{F}$. En projection sur l'axe Ox avec $\vec{T} = -k(x - \ell_0)\vec{u}_x$:

$$m\ddot{x} = -k(x - \ell_0) - \lambda\dot{x} + F$$

soit

$$\ddot{x} + \frac{\omega_0}{Q}\dot{x} + \omega_0^2 x = \omega_0^2 \ell_0 + \frac{F}{m}$$

avec $\omega_0 = \sqrt{k/m}$ et $Q = m\omega_0/\lambda$.

2. On écrit l'équation caractéristique : $r^2 + \frac{\omega_0}{Q}r + \omega_0^2 = 0$. Pour avoir une solution pseudo-périodique,

il faut que le discriminant $\Delta = \omega_0^2 \left(\frac{1}{Q^2} - 4 \right)$ soit négatif donc $Q > 1/2$ ou $\lambda < 2\sqrt{km}$.

Les solutions de l'équation caractéristique sont alors :

$$r_1 = -\frac{\omega_0}{2Q} + \frac{i\omega_0}{2}\sqrt{4 - \frac{1}{Q^2}} \quad \text{et} \quad r_2 = -\frac{\omega_0}{2Q} - \frac{i\omega_0}{2}\sqrt{4 - \frac{1}{Q^2}}$$

La solution particulière constante est $x_p = \ell_0 + \frac{F}{m\omega_0^2}$, d'où :

$$x(t) = A \exp\left(-\frac{\omega_0 t}{2Q}\right) \cos(\Omega t + \varphi) + \ell_0 + \frac{F}{m\omega_0^2}$$

avec $\Omega = \frac{\omega_0}{2}\sqrt{4 - \frac{1}{Q^2}}$. On détermine les constantes A et φ avec les conditions initiales. À $t = 0$, le point M est en $x = \ell_0$ et sa vitesse est nulle :

$$\begin{cases} x(0) = A \cos \varphi + \ell_0 + \frac{F}{m\omega_0^2} = \ell_0 \\ x'(0) = A \left(-\Omega \sin \varphi - \frac{\omega_0}{2Q} \cos \varphi \right) = 0 \end{cases} \Rightarrow \begin{cases} A = -\frac{F}{m\omega_0^2 \cos \varphi} \\ \tan \varphi = -\frac{\omega_0}{2\Omega Q} \end{cases}$$

On peut choisir $\cos \varphi > 0$ donc $\varphi \in]-\frac{\pi}{2}, 0[$ puisque $\tan \varphi < 0$. On peut alors écrire

$$\frac{1}{\cos \varphi} = \sqrt{1 + \tan^2 \varphi} \text{ soit}$$

$$A = -\frac{F}{m\omega_0^2} \sqrt{1 + \left(\frac{\omega_0}{2\Omega Q} \right)^2} \quad \text{et} \quad \varphi = -\arctan \frac{F}{m\omega_0^2 \cos \varphi}$$

3. La pseudo-période est $T = 2\pi/\Omega$ donc :

$$\begin{aligned} x(T) &= A \exp\left(-\frac{2\pi\omega_0}{2Q\Omega}\right) \cos(\Omega T + \varphi) + \ell_0 + \frac{F}{m\omega_0^2} \\ &= A \exp\left(-\frac{2\pi}{\sqrt{4Q^2-1}}\right) \cos \varphi + \ell_0 + \frac{F}{m\omega_0^2} \end{aligned}$$

puis avec l'expression de $A \cos \varphi$:

$$x(T) = \frac{F}{m\omega_0^2} \left(1 - \exp\left(-\frac{2\pi}{\sqrt{4Q^2-1}}\right)\right) + \ell_0$$

A.3 1. Le référentiel du laboratoire est galiléen. Les deux systèmes étudiés sont les points A et B soumis aux forces :

	Poids $m\vec{g}$		Poids $m\vec{g}$
Pour A	Réaction du support $\vec{R}_n$	Pour B	Réaction du support $\vec{R}_n$
	Tension du ressort (1) $\vec{T}_1$		Tension du ressort (2) $\vec{T}_2$
	Tension du ressort (2) $\vec{T}_2$		Tension du ressort (3) $\vec{T}_3$

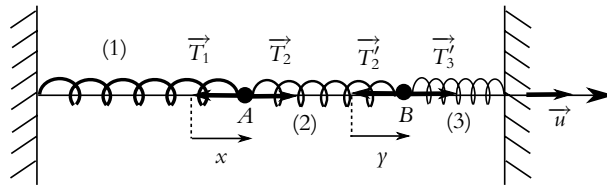


Figure S7.3

Les poids et les réactions verticales n'ont pas été représentés sur la figure (S7.3), on ne s'intéresse ici qu'aux forces horizontales. On définit le vecteur unitaire $\vec{u}$ sur l'axe. On note ℓ_i la longueur du ressort (i).

Pour déterminer l'expression des tensions, on utilise le vecteur $\vec{u}_{\text{ext}}$. Pour le point A , on peut écrire $\vec{u}_{\text{ext}} = \vec{u}$ pour le ressort (1) et $\vec{u}_{\text{ext}} = -\vec{u}$ pour le ressort (2). En ce qui concerne le point B , on a $\vec{u}_{\text{ext}} = \vec{u}$ pour le ressort (2) et $\vec{u}_{\text{ext}} = -\vec{u}$ pour le ressort (3).

Les relations fondamentales à l'équilibre, pour chaque point projetées sur le vecteur $\vec{u}$ donnent :

$$0 = -k(\ell_e - \ell_0) + K(L_e - L_0) \text{ pour } A \text{ et } 0 = +k(\ell_e - \ell_0) - K(L_e - L_0) \text{ pour } B$$

2. Pour établir les équations du mouvement, il faut tout d'abord déterminer les expressions des longueurs des ressorts :

$$\ell_1 = \ell_e + x \quad \ell_2 = L_e - x + y \quad \ell_3 = \ell_e - y$$

On peut alors écrire les relations fondamentales en projection sur $\vec{u}$:

$$\begin{cases} m\ddot{x} = -k(\ell_1 - \ell_0) + K(\ell_2 - L_0) = -(k+K)x + Ky - k(\ell_e - \ell_0) + K(L_e - L_0) \\ m\ddot{y} = +k(\ell_3 - \ell_0) - K(\ell_2 - L_0) = -(k+K)y + Kx + k(\ell_e - \ell_0) - K(L_e - L_0) \end{cases}$$

On peut simplifier les équations précédentes en reportant les relations à l'équilibre. On divise ensuite chaque équation par m , il vient :

$$\begin{cases} \ddot{x} = -\frac{(k+K)}{m}x + \frac{K}{m}y \\ \ddot{y} = -\frac{(k+K)}{m}y + \frac{K}{m}x \end{cases}$$

3. Les deux équations sont dites couplées car elles dépendent toutes les deux des variables x et y . La méthode proposée par l'énoncé pour les résoudre est de faire apparaître les variables $S = x + y$ et

$D = x - y$, pour cela on va effectuer la somme et la différence entre les deux expressions :

$$\begin{cases} \ddot{x} + \ddot{y} = -\frac{k+2K}{m}(x+y) \\ \ddot{x} - \ddot{y} = -\frac{k}{m}(x-y) \end{cases} \Rightarrow \begin{cases} \ddot{S} + \frac{k+2K}{m}S = 0 \\ \ddot{D} + \frac{k}{m}D = 0 \end{cases} \Rightarrow \begin{cases} \ddot{S} + \omega_1^2 S = 0 \\ \ddot{D} + \omega_0^2 D = 0 \end{cases}$$

Le système d'équations précédent nous permet d'écrire :

$$S = A_1 \cos \omega_1 t + A_2 \sin \omega_1 t \quad \text{et} \quad D = B_1 \cos \omega_0 t + B_2 \sin \omega_0 t$$

4. Les conditions initiales sont :

$$\begin{cases} x(0) = 0 \quad \text{et} \quad y(0) = b \\ \dot{x}(0) = 0 \quad \text{et} \quad \dot{y}(0) = 0 \end{cases} \Rightarrow \begin{cases} S(0) = A_1 = b \quad \text{et} \quad D(0) = B_1 = -b \\ \dot{S}(0) = -A_2 \omega_1 = 0 \quad \text{et} \quad \dot{D}(0) = -B_2 \omega_0 = 0 \end{cases}$$

D'où finalement :

$$\begin{cases} S = b \cos \omega_1 t \\ D = -b \cos \omega_0 t \end{cases} \Rightarrow \begin{cases} x = \frac{S+D}{2} = \frac{b}{2}(-\cos \omega_0 t + \cos \omega_1 t) \\ y = \frac{S-D}{2} = \frac{b}{2}(\cos \omega_0 t + \cos \omega_1 t) \end{cases}$$

A.4 1. On utilise la définition de l'énergie potentielle ou on utilise le gradient. On en déduit $\vec{f} = -k_x x \vec{u}_x - k_y y \vec{u}_y - k_z z \vec{u}_z$.

2. On étudie le système M dans le référentiel terrestre galiléen. Le système est soumis à la force $\vec{f}$. Le principe fondamental de la dynamique s'écrit $m \vec{a} = \vec{f}$ dont la projection sur chaque axe donne une équation du type $m\ddot{x} + k_x x = 0$. On en déduit $x = X \cos \left(\sqrt{\frac{k_x}{m}} t + \varphi_x \right)$, $y = Y \cos \left(\sqrt{\frac{k_y}{m}} t + \varphi_y \right)$ et $z = Z \cos \left(\sqrt{\frac{k_z}{m}} t + \varphi_z \right)$.

3. L'énergie mécanique s'écrit

$$Em = Ec + Ep = \frac{1}{2} m (\dot{x}^2 + \dot{y}^2 + \dot{z}^2) + \frac{1}{2} (k_x x^2 + k_y y^2 + k_z z^2)$$

$$\text{soit } Em = \left(\frac{1}{2} m \dot{x}^2 + \frac{1}{2} k_x x^2 \right) + \left(\frac{1}{2} m \dot{y}^2 + \frac{1}{2} k_y y^2 \right) + \left(\frac{1}{2} m \dot{z}^2 + \frac{1}{2} k_z z^2 \right).$$

4. Si $k_x = k_y = k_z = k$, $\vec{f} = -k \overrightarrow{OM}$ et le principe fondamental de la dynamique s'écrit $m \frac{d\overrightarrow{OM}}{dt} + k \overrightarrow{OM} = \vec{0}$.

$$\text{Une résolution vectorielle donne } \overrightarrow{OM} = \overrightarrow{OM}_0 \cos \left(\sqrt{\frac{k}{m}} t \right) + \vec{v}_0 \sqrt{\frac{m}{k}} \sin \left(\sqrt{\frac{k}{m}} t \right).$$

5. On a donc un mouvement dans le plan défini par O , M_0 et $\vec{v}_0$.

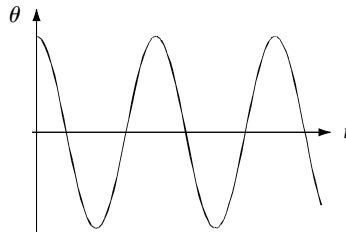
6. Par projection dans ce plan, on obtient $x = x_0 \cos \left(\sqrt{\frac{k}{m}} t \right) + v_{0,x} \sqrt{\frac{m}{k}} \sin \left(\sqrt{\frac{k}{m}} t \right)$ ainsi que

$$y = y_0 \cos \left(\sqrt{\frac{k}{m}} t \right) + v_{0,y} \sqrt{\frac{m}{k}} \sin \left(\sqrt{\frac{k}{m}} t \right). \text{ On « élimine » le temps en utilisant la relation}$$

$$\sin^2 a + \cos^2 a = 1 \text{ valable pour toute valeur de } a, \text{ on en déduit } \frac{x^2}{A^2} + \frac{y^2}{B^2} + \frac{xy}{C} = 1 \text{ avec}$$

$$A^2 = \frac{(v_{0,y} x_0 - v_{0,x} y_0)^2}{v_{0,y}^2 + \frac{k}{m} y_0^2}, \quad B^2 = \frac{(v_{0,x} y_0 - v_{0,y} x_0)^2}{v_{0,x}^2 + \frac{k}{m} x_0^2} \quad \text{et} \quad C = \frac{(v_{0,x} y_0 - v_{0,y} x_0)^2}{v_{0,x} v_{0,y} + \frac{k}{m} x_0 y_0}.$$

- B.1** 1. On étudie le point A de masse m dans le référentiel terrestre galiléen. Ce point est soumis à son poids et à la tension du fil $\vec{T}$.
2. Le principe fondamental de la dynamique projeté sur $\vec{u}_\theta$ en coordonnées polaires donne $mL\ddot{\theta} = -mg \sin \theta$. L'équation du mouvement est donc $\ddot{\theta} + \frac{g}{L} \sin \theta = 0$.
3. Le théorème de l'énergie cinétique s'écrit $dEc = \delta W$ avec $\delta W(\vec{T}) = 0$ puisque la force est perpendiculaire au mouvement et $\delta W(m\vec{g}) = -mgL \sin \theta d\theta$. On en déduit $dEc = mL^2 \dot{\theta} d\theta = -mgL \sin \theta d\theta$. On retrouve donc l'équation du mouvement puisque pour avoir un mouvement, il faut $\dot{\theta} \neq 0$.
4. On peut déterminer l'expression de l'énergie potentielle par $\delta W(m\vec{g}) = -dEp$. On obtient $Ep = -mgL \cos \theta + K$. Les positions d'équilibre correspondent alors à $dEp = 0$ soit $\sin \theta = 0$. On en déduit les deux positions d'équilibre $\theta = 0$ ou $\theta = \pi$.
5. Leur stabilité dépend du signe de $\frac{d^2 Ep}{d\theta^2} = mgL \cos \theta$. Pour $\theta = 0$, l'expression est positive et l'équilibre stable tandis que pour $\theta = \pi$, le signe est négatif et la position d'équilibre instable.
6. En supposant θ faible, on peut faire l'approximation $\sin \theta \simeq \theta$. On en déduit $\ddot{\theta} + \frac{g}{L} \theta = 0$, ce qui correspond à l'équation du mouvement d'un oscillateur harmonique de pulsation $\omega_0 = \sqrt{\frac{g}{L}}$ et de période $T_0 = 2\pi \sqrt{\frac{L}{g}}$.
7. La solution s'écrit $\theta = A \cos \sqrt{\frac{g}{L}} t + B \sin \sqrt{\frac{g}{L}} t$. La détermination des constantes A et B s'obtient à l'aide des conditions initiales soit $\theta(0) = A = \theta_0$ et $\dot{\theta}(0) = B \sqrt{\frac{g}{L}} = 0$. Finalement $\theta = \theta_0 \cos \sqrt{\frac{g}{L}} t$.
8. La vitesse s'écrit $v = l\dot{\theta}$. Elle est donc maximale si $\dot{\theta} = -\theta_0 \sqrt{\frac{g}{L}}$ est maximal soit $v_{\max} = \theta_0 \sqrt{Lg}$.
9. L'évolution de θ avec t est :



10. On doit maintenant ajouter la force de frottement $-h\vec{v}$ soit dans la projection du principe fondamental de la dynamique sur $\vec{u}_\theta$: $\ddot{\theta} + \frac{h}{m} \dot{\theta} + \frac{g}{L} \sin \theta = 0$.
11. Lorsque θ est faible, on obtient $\ddot{\theta} + \frac{h}{m} \dot{\theta} + \frac{g}{L} \theta = 0$. Son équation caractéristique s'écrit $r^2 + \frac{h}{m} r + \frac{g}{L} = 0$ et son discriminant $\Delta = \left(\frac{h}{m}\right)^2 - 4\frac{g}{L}$. On a un régime pseudo-périodique si $\Delta < 0$ soit $h < 2m\sqrt{\frac{g}{L}}$.
12. La solution s'écrit

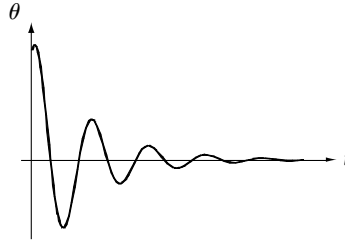
$$\theta = \left(A \cos \sqrt{\frac{g}{L} - \left(\frac{h}{2m}\right)^2} t + B \sin \sqrt{\frac{g}{L} - \left(\frac{h}{2m}\right)^2} t \right) \exp \left(-\frac{h}{2m} t \right).$$

On détermine les constantes A et B avec les conditions initiales $\theta(0) = A = \theta_0$ et

$$\dot{\theta}(0) = \frac{B}{2} \sqrt{\frac{g}{L} - \left(\frac{h}{2m}\right)^2} - \frac{h}{2m} A = 0. \text{ Finalement on a}$$

$$\theta = \left(\cos \sqrt{\frac{g}{L} - \left(\frac{h}{2m}\right)^2} t + h \sqrt{\frac{L}{4gm^2 - Lh^2}} \sin \sqrt{\frac{g}{L} - \left(\frac{h}{2m}\right)^2} t \right) \theta_0 \exp\left(-\frac{h}{2m} t\right)$$

13. L'évolution de θ est donc :



14. On doit ajouter le travail $\delta W(-h\vec{v}) = -hL^2\dot{\theta}d\theta$ et par raisonnement analogue à 3. on retrouve l'équation du mouvement.

B.2 1. On étudie le mouvement de la masse m dans le référentiel du laboratoire considéré comme galiléen. Le système est soumis à son poids $m\vec{g}$ et à la tension $\vec{T}$ du fil. On applique le principe fondamental de la dynamique : $m\vec{a} = \vec{T} + m\vec{g}$ qu'on projette sur la direction $\vec{u}_\theta$ des coordonnées polaires : $-m\ddot{\theta} = -mg \sin \theta$ soit $\ddot{\theta} + \frac{g}{l} \sin \theta = 0$ pour éliminer la force inconnue $\vec{T}$.

2. Si on ne considère que les petites oscillations alors $\sin \theta \simeq \theta$ et on obtient l'équation d'un oscillateur harmonique : $\ddot{\theta} + \frac{g}{l} \theta = 0$ de période $T = 2\pi\sqrt{\frac{l}{g}}$. Cette période ne dépend pas de l'amplitude des oscillations : il y a isochronisme des oscillations.

3. On prend maintenant l'approximation $\sin \theta \simeq \theta - \frac{\theta^3}{6}$ et l'équation du mouvement devient : $\ddot{\theta} + \omega_0^2 \theta - \frac{\omega_0^2}{6} \theta^3 = 0$ en posant $\omega_0 = \sqrt{\frac{g}{l}}$.

4. On cherche des solutions de la forme $\theta = A \cos \omega t$ soit $\dot{\theta} = -\omega A \sin \omega t$ et $\ddot{\theta} = -\omega^2 A \cos \omega t$. En reportant dans l'équation du mouvement, on obtient : $-A\omega^2 \cos \omega t + \omega_0^2 A \cos \omega t - \frac{\omega_0^2}{6} A^3 \cos^3 \omega t = 0$. Or $\cos^3 \alpha = \frac{\cos 3\alpha + 3 \cos \alpha}{4}$ donc en factorisant les cosinus : $A \cos \omega t \left(\omega_0^2 - \omega - \frac{\omega_0^2 A^2}{8} \right) - \frac{\omega_0^2 A^3}{24} \cos 3\omega t = 0$.

Ceci ne sera possible à tout instant t que si $\omega_0^2 - \omega - \frac{\omega_0^2 A^2}{8} = 0$ et $A = 0$. Dans ce cas, il n'y a plus que la solution identiquement nulle, qui est inintéressante physiquement. Il est donc nécessaire d'introduire un mouvement de pulsation triple de la pulsation initiale.

5. On cherche donc une solution sous la forme : $\theta = A (\cos \omega t + \epsilon \cos 3\omega t)$ avec $A \ll 1$ et $\epsilon \ll 1$, ce qui revient à partir de la limite linéaire des petites oscillations et à tenir compte d'un terme correctif par rapport au résultat linéaire.

On a alors $\ddot{\theta} = -A\omega^2 (\cos \omega t + 9\epsilon \cos 3\omega t)$ et $\theta^3 \simeq A^3 \cos^3 \omega t$ en ne gardant que le terme prépondérant du fait que $\epsilon \ll 1$. En linéarisant, on obtient : $\theta^3 \simeq \frac{A^3}{4} (\cos 3\omega t + 3 \cos \omega t)$. Le report dans l'équation du mouvement donne : $A \left(-\omega^2 + \omega_0^2 - \frac{A^2 \omega_0^2}{8} \right) \cos \omega t + A \left(\epsilon \omega_0^2 - 9\epsilon \omega^2 - \frac{A^2 \omega_0^2}{24} \right) \cos 3\omega t = 0$.

Cette équation étant valable quel que soit l'instant t considéré, on a :

$$\begin{cases} \omega_0^2 - \omega^2 - \frac{A^2 \omega_0^2}{8} = 0 \\ \epsilon (\omega_0^2 - 9\omega^2) = \frac{\omega_0^2 A^2}{24} \end{cases}$$

6. De la première condition, on déduit la pulsation fondamentale des oscillations (qui n'est plus la pulsation propre ω_0) : $\omega = \omega_0 \sqrt{1 - \frac{A^2}{8}}$ soit une période $T = \frac{2\pi}{\omega_0 \sqrt{1 - \frac{A^2}{8}}} = T_0 \left(1 - \frac{A^2}{8}\right)^{-\frac{1}{2}}$. Compte tenu du fait que $A \ll 1$, on peut effectuer un développement limité de la période : $T = T_0 \left(1 + \frac{A^2}{16}\right)$. Il s'agit de la formule de Borda. On vient d'établir qu'il n'y a plus isochronisme des oscillations puisque l'expression de la période montre une dépendance avec l'amplitude A . D'autre part, on peut établir l'expression de l'amplitude du mouvement correctif :

$$\epsilon = \frac{\omega_0^2 A^2}{24 \left(\omega_0^2 - 9\omega_0^2 \left(1 - \frac{A^2}{8}\right) \right)} = \frac{A^2}{24 \left(-8 + \frac{9A^2}{8} \right)} \simeq -\frac{A^2}{192}.$$

- B.3** 1. a) Dans le référentiel du laboratoire considéré galiléen, la relation fondamentale de la dynamique appliquée à l'oscillateur M donne en projection sur l'axe Ox :

$$m\ddot{x} = -k(x - \ell_0) - \lambda\dot{x} + F_c \Rightarrow \ddot{x} + \frac{\lambda}{m}\dot{x} + \frac{k}{m}x = \frac{k}{m}\ell_0 + \frac{F_c}{m}$$

On peut donc mettre l'équation sous la forme donnée par l'énoncé avec $\omega_0 = \sqrt{k/m}$, $Q = \sqrt{km}/\lambda$ et $X_0 = \ell_0 + F_c/k$.

- b) Si la solution est pseudo-périodique, alors les solutions de l'équation caractéristique sont :

$$-\frac{\omega_0}{2Q} \pm i\frac{\omega_0}{2}\sqrt{4 - \frac{1}{Q^2}}$$

On peut définir alors la pseudo pulsation $\Omega = \frac{\omega_0}{2}\sqrt{4 - \frac{1}{Q^2}}$ et le temps caractéristique $\tau = \frac{2Q}{\omega_0}$. L'expression de $x(t)$ est alors :

$$x(t) = A \exp(-t/\tau) \cos(\Omega t + \varphi) + X_0$$

où A et φ sont des constantes dépendant des conditions initiales.

2. a) Le mobile passe de part et d'autre de la position finale : il s'agit d'un mouvement pseudo périodique.
- b) Au début du mouvement, le point représentatif de l'état du mobile est A et à la fin il s'agit du point B (on rappelle qu'un portrait de phase est décrit dans le sens horaire). Initialement le mobile est donc en $x = 0$ et sa vitesse est $v_0 = 2 \text{ cm.s}^{-1}$. À la fin M est immobile en $x = 2 \text{ cm}$.
- c) La vitesse maximale atteinte correspond au maximum de la courbe, soit $v = 16,5 \text{ cm.s}^{-1}$ et l'élongation maximale correspond à l'intersection avec l'axe des abscisses la plus à droite soit $x_{\max} \simeq 3,4 \text{ cm}$.
- d) Entre deux croisements avec l'axe des abscisses il se passe une demi pseudo-période. Pour avoir une valeur plus précise, on calcule les intervalles de temps entre chaque date donnée dans l'énoncé ce qui donne quatre valeurs de $T/2$ dont on fait la moyenne. Les quatre valeurs obtenues sont : $0,34 \text{ s}$; $0,32 \text{ s}$; $0,33 \text{ s}$ et $0,32 \text{ s}$; ce qui fait une moyenne de $0,3275 \text{ s}$ soit en arrondissant on a $T = 0,65 \text{ s}$ et $\Omega = 2\pi/T = 9,6 \text{ rad.s}^{-1}$.

- e) En utilisant l'expression générale de $x(t)$ déterminée au début de l'exercice et sachant que $x_B = X_0$ (position finale atteinte lorsque $t \rightarrow \infty$) on a :

$$\begin{aligned}\delta &= \frac{1}{n} \ln \left(\frac{A \exp(-t/\tau) \cos(\Omega t + \varphi)}{A \exp(-(t+nT)/\tau) \cos(\Omega(t+nT) + \varphi)} \right) \\ &= \frac{1}{n} \ln \left(\exp(nT/\tau) \right) = \frac{T}{\tau}\end{aligned}$$

Sur la courbe, on choisit par exemple le premier croisement avec l'axe des abscisses ($x_1 = x_{\max} = 3,4$ cm) et le septième ($x_2 = 2,2$ cm). L'écart temporel entre les deux positions est $3T$ d'où :

$$\delta = \frac{1}{3} \ln \left(\frac{3,4 - 2}{2,2 - 2} \right) = 0,65$$

On en déduit donc $\tau = 1$ s.

- f) On rappelle les expressions de $\Omega = \frac{\omega_0}{2} \sqrt{4 - \frac{1}{Q^2}}$ et $\tau = \frac{2Q}{\omega_0}$. Le produit $\Omega\tau = \sqrt{4Q^2 - 1}$ permet de calculer Q soit $Q = \sqrt{\Omega^2\tau^2 + 1/2}$; l'application numérique donne $Q = 4,8$. On peut alors calculer $\omega_0 = \frac{2Q}{\tau} = 9,65$ rad.s⁻¹.

- g) La constante de raideur du ressort est $k = m\omega_0^2 = 18,6$ N.m⁻¹. Le coefficient de frottement est obtenu avec l'expression de Q : $\lambda = \sqrt{km}/Q = 0,4$ kg.s⁻¹. Pour déterminer F_c , il faut utiliser $X_0 = \ell_0 + F_c/k$ qui correspond à la position de l'attracteur sur le portrait de phase soit $X_0 = 2$ cm. On en déduit $F_c = 0,186$ N.

B.4 1. a) On considère le référentiel terrestre galiléen. Le système est le véhicule modélisé par le point M . On représente sur la figure (S7.4) le système à trois instants : juste avant de franchir le défaut, juste après avoir franchi le défaut (instant pris comme origine des temps), instant t ultérieur. Le défaut de chaussée a été exagéré sur la figure pour plus de clarté. On note Z_i la cote de M avant le défaut.

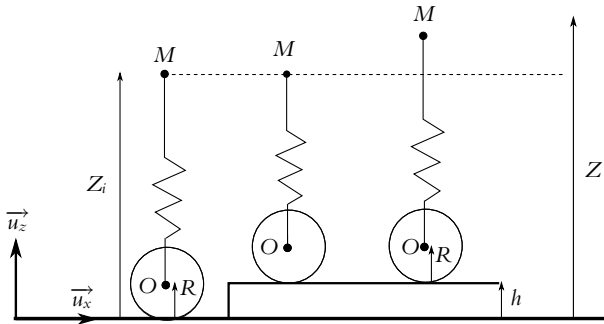


Figure S7.4

Les forces s'exerçant sur M sont : le poids $m\vec{g}$, la tension du ressort $\vec{T} = -k(\ell - \ell_0)\vec{u}_z$ et la force de frottement $\vec{f}_v$. D'après le schéma, la longueur du ressort est $OM = \ell = Z - h - Z_0 = Z - h - R$. La relation fondamentale en projection sur l'axe Oz donne :

$$m\ddot{Z} = -mg - k(Z - h - R - \ell_0) - \lambda\dot{Z}$$

Avant l'obstacle, le point M se déplace à altitude constante, et la longueur du ressort vaut $Z_i - R$, on peut alors écrire l'équation d'équilibre :

$$0 = -mg - k(Z_i - R - \ell_0)$$

L'équation du mouvement devient alors :

$$\ddot{Z} + 2\alpha\dot{Z} + \omega_0^2 Z = \omega_0^2(h + Z_i)$$

avec $\alpha = \lambda/2m = 5 \text{ rad.s}^{-1}$ et $\omega_0^2 = k/m = 50 \text{ rad}^2.\text{s}^{-2}$. Le discriminant de l'équation caractéristique $r^2 + 2\alpha r + \omega_0^2 = 0$ est $\Delta = 4\alpha^2 - 4\omega_0^2 = -100$. Puisque $\Delta < 0$, le mouvement est pseudo-périodique. Les solutions de l'équation caractéristique sont $\alpha \pm i\sqrt{-\Delta}/2$. La solution particulière de l'équation différentielle est $h + Z_i$. On en déduit l'expression de $Z(t)$:

$$Z(t) = A \exp(\alpha t) \cos(\Omega t + \varphi) + h + Z_i$$

avec $\Omega = \sqrt{-\Delta}/2 = 5 \text{ rad.s}^{-1}$, A et φ sont des constantes à déterminer avec les conditions initiales. À $t = 0$ on a $\dot{Z} = v_0$ et $Z(0) = Z_i$. Le calcul de $\dot{Z}$ donne $\dot{Z} = A \exp(\alpha t)(-\alpha \cos(\Omega t + \varphi) - \Omega \sin(\Omega t + \varphi))$ soit avec $\Omega = \alpha$:

$$\dot{Z} = -A\Omega \exp(\alpha t)(\cos(\Omega t + \varphi) + \sin(\Omega t + \varphi))$$

d'où les conditions initiales :

$$\begin{cases} A \cos \varphi + h + Z_i = Z_i \\ -A\Omega(\cos \varphi + \sin \varphi) = v_0 \end{cases} \Rightarrow \begin{cases} A \cos \varphi = -h \\ A \sin \varphi = -\frac{v_0}{\Omega} + h \end{cases}$$

On en déduit $\tan \varphi = \frac{v_0}{h\Omega} - 1 = 1$ soit $\varphi = \pi/4$ et $A = -\sqrt{2}h$.

- b) Pour déterminer la valeur maximale de l'accélération $\ddot{Z}(t)$ il faut dériver cette expression. On utilise à chaque étape de dérivation le fait que $\Omega = \alpha$. On part de la valeur de $\dot{Z}$ déjà calculée. Il vient :

$$\ddot{Z} = A \exp(\alpha t)((\alpha^2 - \Omega^2) \cos(\Omega t + \varphi) + 2\alpha\Omega \sin(\Omega t + \varphi)) = 2A \exp(\alpha t)\Omega^2 \sin(\Omega t + \varphi)$$

puis :

$$\begin{aligned} \ddot{Z} &= 2A\Omega^2 \exp(\alpha t)(-\alpha \sin(\Omega t + \varphi) + \Omega \cos(\Omega t + \varphi)) \\ &= 2A\Omega^3 \exp(\alpha t)(-\sin(\Omega t + \varphi) + \cos(\Omega t + \varphi)) \end{aligned}$$

$\ddot{Z} = 0$ pour $\tan(\Omega t + \varphi) = 1$. Le premier instant pour lequel $\ddot{Z}$ est maximale vérifie soit $\Omega t + \varphi = \pi/4$ soit $t = 0$ puisque $\varphi = \pi/4$.

On en déduit $\ddot{Z}_{\max} = 2A\Omega^2 \sin \varphi = 2,5 \text{ m.s}^{-2}$.

2. On se place dans les conditions critiques en rendant le discriminant nul soit $\Delta = 0$ entraîne $\alpha = \omega_0$ soit $m = \frac{\lambda^2}{4k}$. L'application numérique donne $m = 250 \text{ kg}$. Il faut énormément alléger le véhicule.

B.5 1. On ajoute à l'équation du mouvement d'un oscillateur harmonique la force de frottement solide qui vérifie les conditions rappelées dans l'énoncé. La réaction normale R_N vaut en appliquant le principe fondamental de la dynamique sur la verticale : $R_N = mg$.

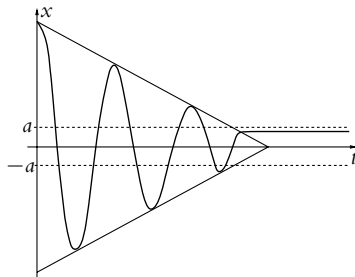
Si $x_0 \leq \frac{fmg}{k} = a$ alors la masse reste en équilibre : la force de rappel a un module inférieur à celui de la force de frottement statique.

2. Si $x_0 > \frac{fmg}{k}$, x diminue dans un premier temps donc $\dot{x} < 0$ et la force de frottement s'écrit : $R_T = +fmg$. L'équation du mouvement est : $m\ddot{x} = -kx + R_T$ soit $\ddot{x} + \frac{k}{m}x = fg$ dont la solution compte tenu des conditions initiales est : $x(t) = a + (x_0 - a) \cos(\omega_0 t)$ en posant $\omega_0 = \sqrt{\frac{k}{m}}$. Ceci est valable tant que $\dot{x} < 0$ c'est-à-dire jusqu'à l'instant $t_1 = \frac{\pi}{\omega_0}$. On note $x(t_1) = 2a - x_0 = x_1$. On recommence alors la même analyse avec, comme nouvelle condition initiale, $x = x_1$ et $\dot{x} = 0$. Si $|x_1| \leq a$ c'est-à-dire si $a \leq x_0 \leq 3a$ alors la masse reste en équilibre.

3. Sinon (à savoir si $x_0 > 3a$, le cas $x_0 < a$ est exclu si on a eu la première phase du mouvement) x augmente et $\dot{x} > 0$ donc $R_T = -mg$ et l'équation du mouvement devient : $\ddot{x} + \frac{k}{m}x = -kfg$. La solution, compte tenu des conditions initiales, s'écrit : $x = -a + (3a - x_0) \cos(\omega_0(t - t_1))$. Ceci est valable tant que $\dot{x} > 0$ c'est-à-dire jusqu'à l'instant $t_2 = \frac{2\pi}{\omega_0}$. On note $x(t_2) = x_0 - 4a = x_2$.

On retrouve le premier cas avec comme condition initiale $x = x_2$ et $\dot{x} = 0$. Il y aura mouvement si $x_0 > 5a$ et alors : $x = a + (x_0 - 5a) \cos(\omega_0(t - t_2))$ jusqu'à $t_3 = \frac{3\pi}{\omega_0}$ avec $x(t_3) = 6a - x_0$.

On obtient l'allure suivante pour $x_0 > 11a$:



Les amplitudes successives des oscillations sont donc $x_0 - \frac{fmg}{k}$, $x_0 - 3\frac{fmg}{k}$, $x_0 - 5\frac{fmg}{k}$, etc. Les extrema suivent une progression arithmétique de raison $2\frac{fmg}{k}$, ce qui correspond à une décroissance linéaire de l'amplitude des oscillations et non exponentielle comme c'était le cas pour un frottement fluide, dans le cas du régime pseudopériodique.

- B.6** 1. La quantité $\frac{7\ell^3 F}{Ed^4}$ est égale à Y donc homogène à une longueur. On effectue une analyse dimensionnelle sachant qu'une force est homogène à $\frac{ML}{T^2}$. Ainsi :

$$\frac{7\ell^3 F}{Ed^4} \sim L \Rightarrow E \sim \frac{L^3}{L^4} \frac{ML}{T^2} \frac{1}{L} \sim \frac{M}{LT^2}$$

E est homogène à une force surfacique (ou à une pression)

- On considère l'extrémité de la fibre comme un point matériel. Les deux forces s'exerçant sur ce point sont F et la tension T du ressort équivalent cherché qui est vers le haut. L'équilibre de l'extrémité donne la relation entre les forces : $T = -F$ or $Y = \frac{7\ell^3 F}{Ed^4}$ soit $T = -\frac{Ed^4}{7\ell^3} Y = -kY$. La fibre est équivalente à un ressort de longueur à vide nulle et constante de raideur $k = \frac{Ed^4}{7\ell^3}$
- L'application numérique donne $k = 2,9 \cdot 10^{-4} \text{ N.m}^{-1}$.
- Soit E_p l'énergie potentielle recherchée ; on peut écrire $dE_p = -\delta W$ où le travail élémentaire de la force est $\delta W = T d\ell$ soit

$$dE_p = -\delta W = -(-k\ell d\ell) \Rightarrow E_p = \frac{1}{2}k\ell^2 + cte$$

On choisit la constante nulle. Dans le cas présent, on a donc :

$$E_p = \frac{1}{2} \frac{Ed^4}{7\ell^3} Y^2$$

- L'énergie mécanique est la somme de l'énergie potentielle trouvée précédemment et de l'énergie cinétique soit :

$$E_m = \frac{Ed^4}{14\ell^3} Y^2 + \rho\ell d^2 \left(\frac{dY}{dt} \right)^2$$

- Il n'y pas d'autres forces à prendre en compte que la tension du ressort équivalent, donc si on applique le théorème de l'énergie cinétique, il vient : $dE_c = \delta W$ or $\delta W = -dE_p$, d'où $dE_c = -dE_p$ soit $dE_m = dE_c + dE_p = 0$. L'énergie mécanique est donc constante.

Pour obtenir l'équation du mouvement, on dérive E_m par rapport au temps :

$$\frac{dE_m}{dt} = 0 \Rightarrow \dot{Y} \left(\frac{Ed^4}{7\ell^3} Y + 2\rho\ell d^2 \ddot{Y} \right) = 0$$

On exclut l'immobilité soit $\dot{Y} = 0 \forall t$, d'où l'équation :

$$\ddot{Y} + \frac{Ed^2}{14\rho\ell^4} Y = 0$$

qui est l'équation d'un oscillateur harmonique.

7. La pulsation propre de l'oscillateur précédent est $\omega_0 = \sqrt{\frac{Ed^2}{14\rho\ell^4}}$ et la fréquence propre $f_0 = \omega_0/2\pi$.

8. L'application numérique donne : $\omega_0 = 288 \text{ rad.s}^{-1}$ et $f_0 = 45,9 \text{ Hz}$.

B.7 Le référentiel terrestre est supposé galiléen. Le système étudié est le point G représentant le canon.

1. Les forces s'exerçant sur le canon sont : son poids (vertical), une réaction de support (verticale) et la tension $\vec{T}$ du ressort qui est la seule force horizontale. Au repos $\vec{T} = 0$ donc la longueur ℓ du ressort est égale à L_0 .
2. Le déplacement de G est horizontal, donc la seule force qui travaille est la tension ; cette force est conservative donc l'énergie mécanique est constante. La longueur du ressort étant égale à x , l'expression de l'énergie mécanique est :

$$E_m = \frac{1}{2} M \dot{x}^2 + \frac{1}{2} k_1 (x - L_0)^2$$

À l'instant initial, juste après le départ de l'obus, le canon n'a pas encore bougé donc $x = L_0$ et d'après l'énoncé $v = v_c$. Lorsque G atteint la distance de recul d sa vitesse est nulle et $x = L_0 - d$. On écrit l'égalité de l'énergie mécanique dans ces deux cas :

$$\frac{1}{2} M v_c^2 = \frac{1}{2} k_1 d^2 \Rightarrow d = \sqrt{\frac{M}{k_1}} v_c = \frac{m}{\sqrt{k_1 M}} v_0 \Rightarrow k_1 = \frac{m^2 v_0^2}{d^2 M}$$

L'application numérique donne $k = 1800 \text{ N.m}^{-1}$.

3. Si on applique la relation fondamentale de la dynamique en projection sur Ox, il vient : $M\ddot{x} = -k_1(x - L_0)$ soit $\ddot{x} + \omega_0^2 x = \omega_0^2 L_0$, avec $\omega_0 = \sqrt{k_1/M}$. La solution est $x = A \cos(\omega_0 t + \varphi) + L_0$. À $t = 0$, $x = L_0$ ce qui donne $A \cos \varphi = 0$ donc on choisit par exemple $\varphi = -\pi/2$. Ainsi on peut écrire $x = A \sin \omega_0 t + L_0$. Pour la vitesse, on calcule la dérivée soit $\dot{x} = A \omega_0 \cos \omega_0 t$. À $t = 0$, $\dot{x} = -mv_0/M$ d'où l'expression de A et finalement :

$$x(t) = -\frac{mv_0}{M\omega_0} \sin \omega_0 t + L_0 = -\frac{mv_0}{\sqrt{k_1 M}} \sin \omega_0 t + L_0$$

La valeur absolue de l'amplitude du sinus est égale à d ce qui donne la même valeur que précédemment.

4. Comme le montre la solution de l'équation trouvée dans la question précédente, si l'on utilise un ressort seul, le canon va osciller et donc après le recul, il va repartir vers l'avant. L'amplitude va diminuer petit à petit à cause des frottements inéluctables mais le temps d'immobilisation sera important. On a donc intérêt à ajouter une force de frottement visqueux.
5. La variation d'énergie mécanique est maintenant égale à l'énergie absorbée par le dispositif de freinage, soit $\Delta E_m = -E_a$. Or initialement $E_m = E_c = Mv_c^2/2$ et finalement $E_m = E_p = k_2 d^2/2$, d'où :

$$\frac{1}{2} k_2 d^2 - \frac{1}{2} M v_c^2 = -E_a \Rightarrow k_2 = \frac{1}{d^2} (M v_c^2 - 2E_a) = \frac{1}{d^2} \left(\frac{m^2}{M} v_0^2 - 2E_a \right)$$

L'application numérique donne $k_2 = 244 \text{ N.m}^{-1}$. La pulsation propre est $\omega_0 = \sqrt{k_2/M}$ ce qui donne $0,55 \text{ rad.s}^{-1}$.

6. La relation fondamentale en projection sur Ox donne :

$$\ddot{x} + \frac{\lambda}{M} \dot{x} + \omega_0^2 x = \omega_0^2 L_0$$

le discriminant de l'équation caractéristique $\Delta = \left(\frac{\lambda}{M}\right)^2 - 4\omega_0^2$ doit être nul pour le régime critique d'où $\lambda = 2M\omega_0$. L'application numérique donne $\lambda = 884 \text{ kg.s}^{-1}$.

7. La solution de l'équation est alors :

$$x(t) = \exp\left(-\frac{\lambda t}{2M}\right) (At + B) + L_0$$

À $t = 0$, $x = L_0$ d'où $B = 0$. On calcule la vitesse :

$$\dot{x} = A \exp\left(-\frac{\lambda t}{2M}\right) \left(1 - \frac{\lambda}{2M}t\right)$$

soit à $t = 0$, $\dot{x} = A = v_c$. Finalement on trouve :

$$x(t) = v_c t \exp\left(-\frac{\lambda t}{2M}\right) + L_0 = -\frac{m}{M} v_0 t \exp\left(-\frac{\lambda t}{2M}\right) + L_0$$

Le recul est maximal lorsque la vitesse s'annule soit pour $t_m = 2M/\lambda = 1,8$ s. On peut alors écrire la valeur de x à t_m :

$$x(t_m) = -\frac{m}{M} v_0 \frac{2M}{\lambda} e^{-1} + L_0 = -\frac{2mv_0}{\lambda e} + L_0 = L_0 - d \Rightarrow d = \frac{2mv_0}{\lambda e} = 1 \text{ m}$$

On retrouve la valeur de d précédente.

Partie III – Optique I

Chapitre 9

A.1 L'observateur a l'impression que l'objet est dans la direction du rayon qu'il reçoit. On voit sur la figure (S9.1) que le rayon provenant de l'objet s'est éloigné de la normale au dioptre (air-eau) puisque $n > 1$. On note A la position réelle de l'objet et B sa position apparente.

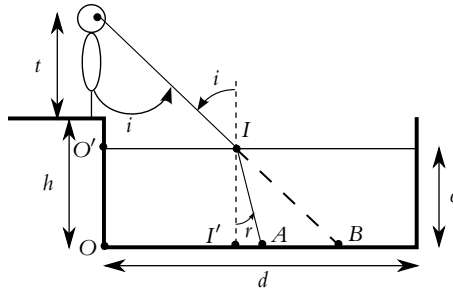


Figure S9.1

- Dans le triangle TOB : $d = OB = (h + t) \tan i$ d'où $i = 16,9^\circ$.
- Dans le triangle $TO'I$: $O'I = (h - e + t) \tan i$.
- Dans le triangle $II'A$: $AI' = e \tan r$.
- Loi de Descartes en I : $\sin i = n \sin r$.

D'autre part, les angles sont inférieurs à $\pi/2$ donc :

$$\cos r = \sqrt{1 - \sin^2 r} = \sqrt{1 - \left(\frac{\sin i}{n}\right)^2}$$

On en déduit que :

$$OA = O'I + I'A = e \tan r + (h - e + t) \tan i$$

donc :

$$OA = e \frac{\sin i}{n} \frac{1}{\sqrt{1 - \left(\frac{\sin i}{n}\right)^2}} + (h - e + t) \tan i = 0,92 \text{ m}$$

- A.2** 1. Il s'agit de calculer l'angle d'incidence i (figure S9.2). Dans le triangle IOH , on a $\sin i = \frac{h}{R} \leq \frac{1}{2}$, or $\text{Arcsin}(0,5) = 0,52 \text{ rad}$. On commet donc au maximum une erreur de 2 sur 50 soit 4 pour cent. Dans la suite, on prend donc $h = Ri$.

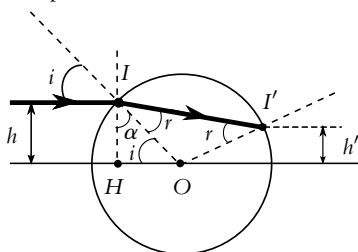


Figure S9.2

2. Le triangle IOI' est isocèle donc $II' = 2R \cos r \simeq 2R$ dans l'approximation de Gauss. D'autre part, $h - h' = II' \cos \alpha$ avec $\alpha = r + \pi/2 - i$ d'où :

$$h' = h - II' \cos(r + \frac{\pi}{2} - i) = h + 2R \sin(r - i) \simeq h + 2R(r - i)$$

En appliquant la loi de Descartes en I dans l'approximation de Gauss : $\sin i = n \sin r$ soit $i = nr$, on en déduit :

$$h' \simeq h + 2Ri \left(\frac{1}{n} - 1 \right) \simeq h + 2h \left(\frac{1}{n} - 1 \right) = h \left(\frac{2}{n} - 1 \right)$$

On obtient $h' = 0$ pour $n = 2$ et pour tous les rayons considérés.

En fait les verres d'indice 2 n'existent pas mais on peut s'en rapprocher.

3. Si $h' = 0$, l'angle d'incidence sur le deuxième dioptré est r par rapport à l'axe (figure S9.3). Le rayon réfléchi repart symétriquement, et la loi de Descartes en J est la même qu'en I . Le rayon renvoyé par la bille à la même direction que le rayon incident et repart donc dans la direction de l'automobiliste.

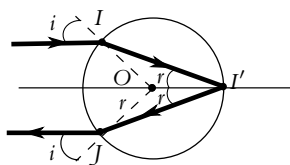


Figure S9.3

- A.3** 1. a) Il peut y avoir réflexion totale en I , car le rayon passe du verre à l'air qui est moins réfringent. On calcule l'angle de réflexion limite en I : $R_{\text{lim},v} = \text{Arcsin} \frac{1}{n_v} = 41,8^\circ$ (on a pris 1 pour l'indice de l'air). Il y a réflexion totale en I si $i > R_{\text{lim},v}$ ce qui est le cas pour $i = 60^\circ$.
- b) On applique la loi de réflexion de Descartes en I ; le rayon est réfléchi symétriquement par rapport à la normale : le détecteur D doit être placé symétriquement à E par rapport à la normale.
2. a) Les deux dioptrés sont maintenant verre-eau et eau-air. Dans les deux cas, il peut y avoir réflexion totale. On calcule les deux angles de réflexion totale :
- Pour le dioptré verre-eau : $n_v \sin R_{\text{lim},1} = n_e \sin \frac{\pi}{2}$ soit $R_{\text{lim},1} = 62,5^\circ$.
 - Pour le dioptré eau-air : $n_e \sin R_{\text{lim},2} = \sin \frac{\pi}{2}$ soit $R_{\text{lim},2} = 48,75^\circ$.

L'angle i est inférieur à $R_{\text{lim},1}$, il y a donc réfraction sur le premier dioptre. On calcule l'angle r après réfraction :

$$n_v \sin i = n_e \sin r \Rightarrow r = 77,6^\circ$$

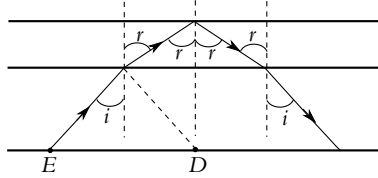


Figure S9.4

L'angle r est aussi l'angle d'incidence sur le deuxième dioptre, il y a donc réflexion totale sur ce dioptre puisque $r > R_{\text{lim},2}$. On obtient la marche du rayon lumineux représenté sur la figure (S9.4).

- b) La figure (S9.4) montre que le rayon lumineux ne tombe plus sur le détecteur lorsqu'il y a de l'eau sur le pare-brise. Un système de commande relié au détecteur peut alors déclencher les essuie-glaces.

A.4 L'hypothénuse étant argentée se comporte comme un miroir plan ; le système est donc équivalent à une lame à faces parallèles en remplaçant le triangle rectangle isocèle par un carré. La loi de Descartes s'écrit $n \sin r = \sin i$ avec $i = 45^\circ$. La distance b parcourue dans le prisme vérifie $\cos r = \frac{a}{b}$ en notant a le côté du carré. Le calcul de $x + y$ s'obtient à partir de l'angle α entre la direction incidente et le parcours de la lumière dans le prisme. On a $i + 90^\circ + r + \alpha = 180^\circ$ soit $\alpha = 90^\circ - i - r$ et $\sin \alpha = \frac{x + y}{b}$. On en déduit $x + y = b \sin \alpha = \frac{a}{\cos r} \sin \alpha = \frac{a}{\cos r} \sin (45^\circ - r)$. En développant, on obtient $x + y = \frac{a}{\sqrt{2}} (1 - \tan r)$. Ce résultat est indépendant de la position d'arrivée du rayon sur la face du prisme.

A.5 On note i l'angle d'incidence, i_1 l'angle de réflexion et i_2 l'angle de réfraction. Les lois de Descartes s'écrivent $i = -i_1$ et $n \sin i = n' \sin i_2$ et on cherche $i_2 - i_1 = \frac{\pi}{2}$. En combinant les relations, on obtient $n \sin (-i_1) = n' \sin i_2$ avec $-i_1 = \frac{\pi}{2} - i_2$ soit $n \cos i_2 = n' \sin i_2$. On a donc $\tan i_2 = \frac{n}{n'}$.

Application numérique : $n = 1,00$ et $n' = 1,33$ soit $i_2 = 36^\circ 56'$.

B.1 On raisonne d'après la figure de l'énoncé sur des angles positifs (ils ne sont pas orientés).

1. L'angle d'incidence limite (dans le verre) R_{lim} correspondant à la réflexion totale est obtenu pour un angle d'émergence dans l'air de $\pi/2$. On écrit la loi de Descartes pour le dioptre, soit $n \sin R_{\text{lim}} = 1$ d'où $R_{\text{lim}} = \text{Arcsin} \frac{1}{n} = 41,8^\circ$.
2. a) L'angle d'incidence du rayon (OA) sur les miroirs est noté ϵ (voir figure S9.5). Le rayon réfléchi (AB) est par hypothèse selon l'axe optique de la lentille de projection. L'angle entre la normale aux miroirs et (AB) est α . La loi de Descartes de la réflexion impose alors $\epsilon = \alpha$. Le segment (OA) est la diagonale du rectangle $(ABOC)$, donc $\theta_r = \alpha + \epsilon = 2\alpha$. Il suffit alors d'appliquer la relation de Descartes en O :

$$n \sin \theta_a = \sin \theta_r \Rightarrow \theta_a = \text{Arcsin} \left(\frac{1}{n} \sin 2\alpha \right) = 13,2^\circ$$

- b) La figure (S9.6) permet d'écrire différentes relations entre les angles :

$$\begin{cases} \theta = \frac{\pi}{2} - i \\ \text{triangle (OKE)} : \beta + \theta_a + \frac{\pi}{2} + \theta = \pi \end{cases}$$

On en déduit $\theta_a = i - \beta$.

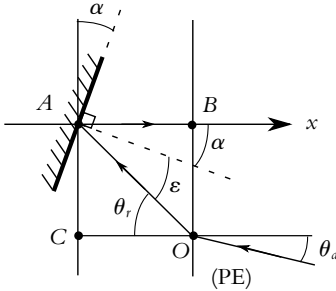


Figure S9.5

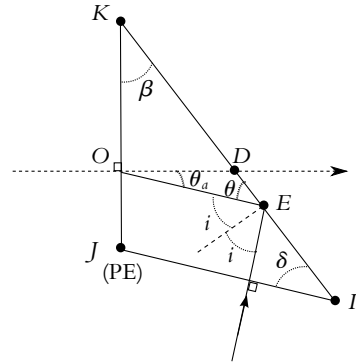


Figure S9.6

- c) En E , il s'agit d'une réflexion verre-air, car l'énoncé signale qu'il y a une mince couche d'air entre les deux prismes. L'angle i doit donc vérifier l'inégalité $i \geq R_{\text{lim}}$, soit pour β :

$$\beta \geq R_{\text{lim}} - \theta_a \text{ ou } \beta \geq \beta_1 \text{ avec } \beta_1 = R_{\text{lim}} - \theta_a = 28,6^\circ$$

- d) Pour qu'il y ait réfraction au point D (interface air-verre), il faut cette fois-ci que l'angle d'incidence soit inférieur à R_{lim} . La normale au dioptre KI en D fait l'angle β avec l'axe OD , et cela correspond à l'angle d'incidence en D . Il faut donc $\beta \leq R_{\text{lim}}$ soit $\beta_2 = R_{\text{lim}} = 41,8^\circ$.

- e) La relation démontrée à la question (b) donne $\beta = i - \theta_a = 31,8^\circ$. Cette valeur est bien comprise entre β_1 et β_2 .

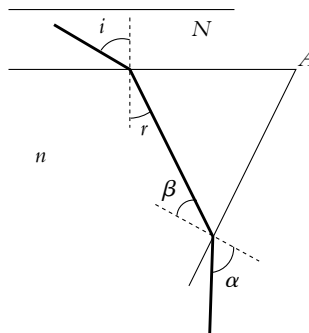
D'autre part, d'après la figure de l'énoncé, on peut écrire : $\delta = \frac{\pi}{2} - \gamma = i$

Il ne peut y avoir de réflexion pour le deuxième prisme car la lumière passe de l'air à un milieu plus réfringent.

- B.2** 1. Les lois de Descartes relatives à la réflexion sont que le rayon réfracté est dans le plan d'incidence défini par le rayon incident et la normale à la surface au point d'arrivée du rayon sur la surface, le rayon réfracté traverse la normale et l'angle de réfraction vérifie $n_1 \sin i_1 = n_2 \sin i_2$ (on indice par 1 le milieu incident et par 2 le milieu réfracté). On a donc un rayon réfracté si $|\sin i_2| \leq 1$ soit $i_1 \leq \text{Arcsin} \frac{n_2}{n_1}$.

2. On a réflexion totale si $n_1 > n_2$ et si $i_1 > \text{Arcsin} \frac{n_2}{n_1}$.

3. Le schéma du montage est :



4. L'air a pour indice 1,00 et le milieu N avec $1 < N$ donc il n'y a pas de réflexion totale.

5. On n'aura pas de réflexion totale entre le matériau d'indice N et le prisme d'indice n si $n > N$.
6. D'après la question précédente, le fait que $N < n$ impose que le rayon soit toujours réfracté. L'angle i varie entre 0 et $\frac{\pi}{2}$ soit par application de la troisième loi de Snell-Descartes r varie entre 0 et $r_t = \text{Arcsin} \frac{N}{n}$.
7. La somme des angles dans un triangle vaut π . En appliquant cette propriété, on obtient $\beta = A - r$. Comme $0 \leq r \leq r_t$, on en déduit $A - r_t \leq \beta \leq A$.
8. Comme $n > 1$, une réflexion totale est possible si $\beta > \text{Arcsin} \frac{1}{n}$. On en déduit :
 - si $A < \text{Arcsin} \frac{1}{n}$, il n'y a pas de réflexion totale car $\beta \leq \text{Arcsin} \frac{1}{n}$,
 - si $A - r_t > \text{Arcsin} \frac{1}{n}$, on a réflexion totale car $\beta \geq \text{Arcsin} \frac{1}{n}$,
 - si $A - r_t < \text{Arcsin} \frac{1}{n} < A$, la réflexion totale n'a lieu que si $\text{Arcsin} \frac{1}{n} < \beta < A$.
9. Finalement on a un rayon émergent si $A < \text{Arcsin} \frac{1}{n}$ ou si $A - r_t < \beta < \text{Arcsin} \frac{1}{n}$ lorsque $A > \text{Arcsin} \frac{1}{n}$.
10. L'application numérique donne $\text{Arcsin} \frac{1}{n} = 35,2^\circ$ donc on est dans le deuxième cas et $A - r_t < \beta < 35,2^\circ$.
11. On a une plage sombre entre 0 et α_t (correspondant à r_t), et une plage claire au-delà.
12. L'application de la troisième loi de Descartes donne

$$N = n \sin \left(\frac{\pi}{2} - r_t \right) = n \sqrt{1 - \left(\frac{\sin^2 \alpha}{n^2} \right)^2} = 1,658$$

pour $\alpha = 30^\circ$.

13. De même, l'application numérique donne $\text{Arcsin} \frac{1}{n} = 38,7^\circ$ et on est toujours dans le deuxième cas soit $A - r_t < \beta < 38,7^\circ$.
14. L'application de la troisième loi de Descartes

$$\sin \alpha = n \sin (61^\circ - r_t) = n \sin 61^\circ \sqrt{1 - \frac{N^2}{n^2}} - n \cos 61^\circ \frac{N}{n}$$

conduit à l'équation $N^2 + 2N \sin \alpha \cos A + \sin^2 \alpha - n^2 \sin^2 A = 0$.

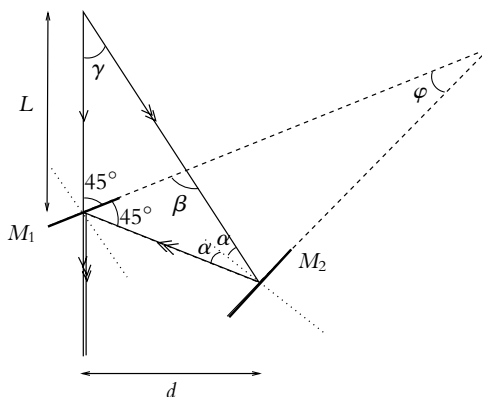
15. L'équation du second degré admet comme discriminant

$$\Delta = 4 \sin^2 \alpha \cos^2 61^\circ - 4 (\sin^2 \alpha - n^2 \sin^2 61^\circ) = 7,623 > 0$$

pour $\alpha = 15^\circ$ et $N = \frac{-2 \sin \alpha \cos 61^\circ \pm \sqrt{\Delta}}{2} = 1,255$.

16. On obtient la même chose car $N = n' \sin r'_t = n \sin r$.
17. On doit considérer la limite entre le liquide et le prisme donc on a le même raisonnement en remplaçant le matériau par le liquide.
18. Si $n' < N$, on n'a pas de réflexion totale à la frontière matériau - liquide mais à l'interface liquide - prisme. Dans ce cas, on mesure l'indice du liquide et non du matériau par le même type de relation.

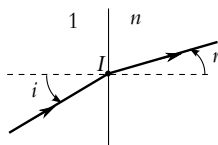
B.3



1. L'application des relations de Descartes donne $\sin i = n \sin r = \sin i'$ soit $i = i'$: les rayons incident et émergent sont parallèles.
2. La somme des angles d'un triangle est égale à π soit en appliquant cette relation aux triangles O_1AO_2 d'une part $\frac{\pi}{4} + \beta + 2\alpha = \pi$ et AO_1B d'autre part $\pi - \beta + \frac{\pi}{2} - \alpha + \varphi = \pi$. Alors en éliminant α entre les deux relations, on en déduit $\beta = \frac{\pi}{4} + 2\varphi$.
3. La somme des angles d'un triangle est égale à π soit pour SO_1A $\frac{\pi}{4} + \gamma + \pi - \beta = \pi$. En tenant compte du résultat de la question précédente, on en déduit $\gamma = 2\varphi$.
4. L'expression de la tangente dans le triangle rectangle donne $\tan 2\varphi = \frac{d}{L}$ dont on déduit $L = \frac{d}{\tan 2\varphi}$.

B.4

1. Un milieu est dispersif est un milieu qui se comporte différemment en fonction de la longueur d'onde (ou de la fréquence) d'une onde électromagnétique. Dans le cas considéré, il s'agit de l'eau qui est un milieu dispersif pour la lumière (comme tous les milieux transparents) car l'indice de réfraction suit la loi de Cauchy : $n = a + \frac{b}{\lambda^2}$, où a et b sont des constantes positives caractéristiques du milieu et λ la longueur d'onde.
2. a) Loi de Descartes pour la réfraction : le rayon réfracté appartient au plan d'incidence et : $\sin i = n \sin r$.



On dérive cette loi par rapport à i sachant que n est indépendant de i :

$$\cos i = n \cos r \frac{dr}{di} \Rightarrow \frac{dr}{di} = \frac{1 \cos i}{n \cos r}$$

Les angles sont dans l'intervalle $[-\frac{\pi}{2}, \frac{\pi}{2}]$, donc le cosinus est positif; on peut utiliser la formule $\cos = \sqrt{1 - \sin^2}$:

$$\frac{dr}{di} = \frac{1}{n} \frac{\sqrt{1 - \sin^2 i}}{\sqrt{1 - \sin^2 r}} = \frac{1}{n} \frac{\sqrt{1 - \sin^2 i}}{\sqrt{1 - \left(\frac{\sin i}{n}\right)^2}}$$

- b) La déviation est l'angle dont a tourné le rayon lumineux par rapport à la direction incidente. Dans le cas de la réfraction (figure S9.7), la déviation D est égale $i - r$, soit en valeur absolue $|D| = |i - r|$.

- c) Dans le cas de la réflexion (figure S9.8), la déviation vaut $\pi - i + r$, si l'on choisit le sens horaire comme positif (dans le cas de la figure $r < 0$ d'où le signe « + »). Avec la loi de Descartes pour la réflexion, $r = -i$ soit $D = \pi - 2i$.

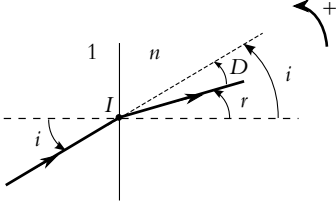


Figure S9.7

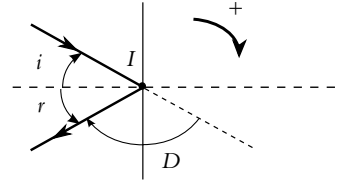


Figure S9.8

3. Les rayons émergents sont parallèles quel que soit l'angle d'incidence i si la déviation D est indépendante de i soit $\frac{dD}{di} = 0$.
4. On considère des angles arithmétiques (positifs)

Cas 1

- a) Le triangle OIJ est isocèle, donc $\alpha = r$. Dans ce cas, la loi de Descartes en J est la même qu'en I et $\beta = i$.
- b) La déviation en I est $i - r$, et celle en J est $\beta - \alpha = i - r$, soit globalement $D_1 = 2(i - r)$.
- c) On calcule la dérivée de D_1 par rapport à i :

$$\frac{dD_1}{di} = 2 - 2\frac{dr}{di} = 2 - 2\frac{1}{n} \frac{\sqrt{1 - \sin^2 i}}{\sqrt{1 - \left(\frac{\sin i}{n}\right)^2}}$$

La dérivée est nulle si $\sqrt{1 - \sin^2 i} = n\sqrt{1 - \left(\frac{\sin i}{n}\right)^2}$. En élevant au carré et en simplifiant, on trouve $n = 1$ ce qui n'a pas d'intérêt. Dans ce cas, il n'est pas possible d'avoir des rayons émergents parallèles (C'est comme dans le cas du prisme).

Cas 2

- a) L'application de la loi de Descartes pour la réflexion en J donne $\beta = \alpha$. D'autre part, les triangles OIJ et OJK sont isocèles donc $\alpha = \beta = \gamma = r$. Alors, l'application de la loi de Descartes en K (identique à celle en I) donne $\delta = i$.
- b) Pour la déviation, on a la somme de trois termes :

$$\begin{cases} i - r \text{ en } I \\ \pi - 2\alpha = \pi - 2r \text{ en } J \\ \delta - \gamma = i - r \text{ en } K \end{cases}$$

Soit globalement :

$$D_2 = \pi + 2i - 4r \quad (1)$$

- c) On dérive D_2 :

$$\frac{dD_2}{di} = 2 - 4\frac{dr}{di} = 2 - 4\frac{1}{n} \frac{\sqrt{1 - \sin^2 i}}{\sqrt{1 - \left(\frac{\sin i}{n}\right)^2}}$$

Cette dérivée s'annule si :

$$\sin^2 i = \frac{4 - n^2}{3} \quad (2)$$

Cas 3

- a) L'application de la loi de Descartes relative à la réflexion en J et K donne : $\gamma = \delta$ et $\alpha = \beta$. De plus, les triangles OIJ , OJK et OKL sont isocèles donc $\alpha = \beta = \gamma = \delta = \varphi = r$. Enfin, l'application de la loi de Descartes en L (identique à celle en I) donne $\xi = i$.

b) Pour la déviation, on a la somme de quatre termes :

$$\begin{cases} i - r \text{ en } I \\ \pi - 2\alpha = \pi - 2r \text{ en } J \\ \pi - 2\gamma = \pi - 2r \text{ en } K \\ \xi - \varphi = i - r \text{ en } L \end{cases}$$

Soit globalement :

$$D_3 = 2\pi + 2i - 6r \quad (3)$$

c) On dérive D_3 :

$$\frac{dD_3}{di} = 2 - 6 \frac{dr}{di} = 2 - 6 \frac{1}{n} \frac{\sqrt{1 - \sin^2 i}}{\sqrt{1 - \left(\frac{\sin i}{n}\right)^2}}$$

Cette dérivée s'annule si :

$$\sin^2 i = \frac{9 - n^2}{8} \quad (4)$$

5. Le soleil étant très bas sur l'horizon, on peut supposer que les rayons incidents sont horizontaux. La figure (S9.9) représente les deux cas possibles. Les angles θ_2 et θ_3 entre les rayons incidents pour l'œil et les rayons provenant du soleil valent :

$$\begin{cases} \theta_2 = \pi - D_2 \\ \theta_3 = D_3 - \pi \end{cases}$$

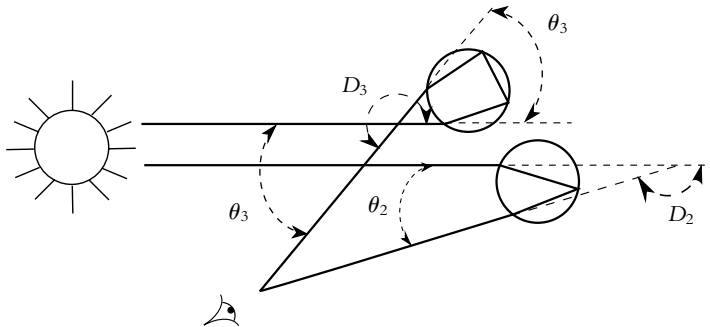


Figure S9.9

6. Pour l'arc primaire, on calcule i avec la relation (2), on en déduit r avec la loi de Descartes puis D_2 avec la relation (1). On procède de même pour D_3 avec les relations (4) et (3).

couleur	i	r	D_2	θ_2	$\theta_2(^{\circ})$
rouge	1.0382	0.7035	2.4039	0.7377	42.2675
violet	1.0250	0.6887	2.4366	0.7050	40.3927
couleur	i	r	D_3	θ_3	$\theta_3(^{\circ})$
rouge	1.2546	0.7948	4.0238	0.8823	50.5494
violet	1.2473	0.7825	4.0830	0.9414	53.9363

L'angle θ_2 est plus grand pour le rouge que pour le violet. Les rayons rouges observés par l'œil (figure S9.10) proviennent donc de gouttes de plus haute altitude, car dans ce cas les rayons violets « passent au-dessus de l'œil », et ce sont les rayons violets des gouttes de plus basse altitude qui entrent dans l'œil. On voit donc le violet en dessous du rouge pour l'arc primaire.

En ce qui concerne, l'arc secondaire, il n'est pas toujours visible car beaucoup moins intense (il y a deux réflexions vitreuses dans la goutte). On remarque que $\theta_3 > \theta_2$, donc les gouttes donnant l'arc secondaire visible par l'œil seront à plus haute altitude que celles donnant l'arc primaire (voir

figure S9.9). L'arc secondaire est donc vu au-dessus de l'arc primaire. D'autre part, les couleurs sont inversées car θ_3 est plus grand pour le violet que pour le rouge. La figure (S9.11) récapitule l'ordre des couleurs.

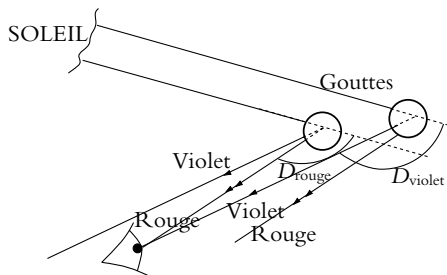


Figure S9.10

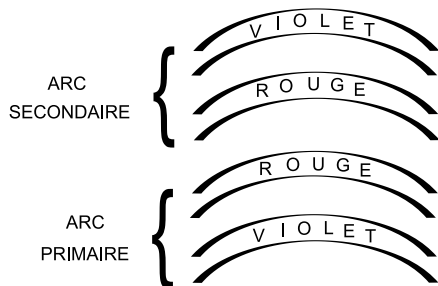


Figure S9.11

Chapitre 10

A.1 1. On note A l'objet et A' l'image. On cherche les images réelles donc positionnées avant le miroir (soit avant son sommet S) : on utilisera donc la formule de conjugaison avec origine au sommet :

$$\frac{1}{\overline{SA'}} + \frac{1}{\overline{SA}} = \frac{2}{\overline{SC}}$$

L'image est réelle si $\frac{1}{\overline{SA'}} < 0$, soit $\frac{1}{\overline{SA}} > \frac{2}{\overline{SC}}$. On sait que pour un miroir concave $\overline{SC} < 0$; il y a donc deux possibilités :

$$\begin{cases} \overline{SA} > 0 \text{ l'inégalité est toujours vérifiée} \\ \overline{SA} < 0 \text{ et } \overline{AS} > \frac{\overline{CS}}{2} \end{cases}$$

Les deux solutions illustrées sur les figures suivantes sont donc : A virtuel ou A réel avant F , puisque le foyer F est à mi-distance de S et C .

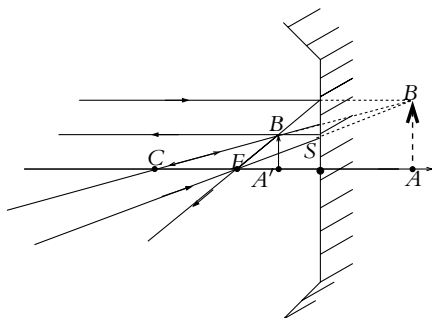


Figure S10.1 A virtuel

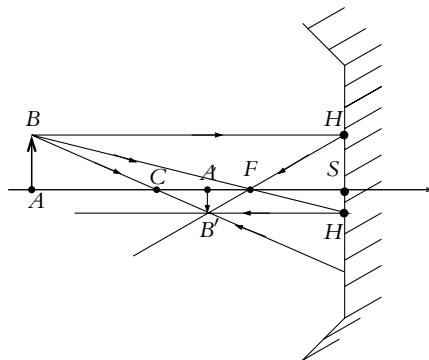


Figure S10.2 A réel

2. Dans le cas d'un miroir convexe, $\overline{SC} > 0$. La seule possibilité est alors $\overline{SA} > 0$ et $\overline{SA} < \frac{\overline{SC}}{2}$: l'objet est virtuel et placé entre F et S comme illustré sur la figure (S10.3).

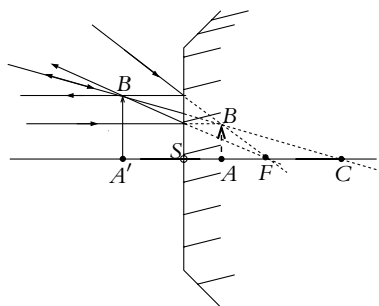
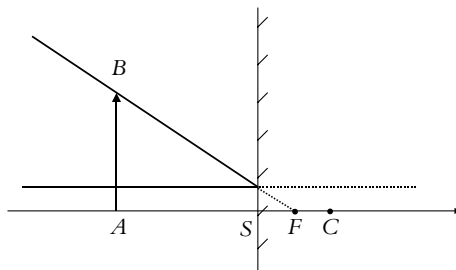


Figure S10.3

A.2 Par le calcul, on utilise la relation de conjugaison $\frac{1}{SA'} + \frac{1}{SA} = \frac{2}{SC}$ et l'expression du grandissement $\gamma = -\frac{SA'}{SA}$. En reportant l'expression de SA' obtenue par le grandissement dans la relation de conjugaison, on en déduit $\overline{SC} = \frac{2SA}{1 - \frac{1}{\gamma}} = 0,30 \text{ m}$.

Par une construction géométrique, on place l'objet à 60 cm devant le miroir puis on trace le rayon passant par B et arrivant sur le miroir à une hauteur 5 fois plus faible. Ce rayon intercepte l'axe optique au foyer objet et ressort parallèlement à l'axe optique. On en déduit la position du centre du miroir à partir du fait que F est le milieu de [SC].



A.3 1. La réponse est (c).

En effet, la formule de conjugaison des miroirs est $\frac{1}{SA'} + \frac{1}{SA} = \frac{2}{SC}$. Le foyer objet F (respectivement image F') correspond à A' (respectivement A) à l'infini soit F et F' au milieu de [SC].

2. La réponse est (a).

En utilisant les rayons arrivant au sommet et repartant symétriquement par rapport à l'axe optique, on exprime la tangente de l'angle d'incidence $\tan i = \frac{AB}{SA} = \frac{A'B'}{A'S}$ dont on déduit $\frac{A'B'}{AB} = -\frac{SA'}{SA} = \frac{1}{5}$. La formule de conjugaison permet alors d'obtenir $V = \frac{1}{SF} = +0,4\delta$.

3. La réponse est (c).

En effet, on a $SF > 0$ donc le miroir est par définition convexe et divergent.

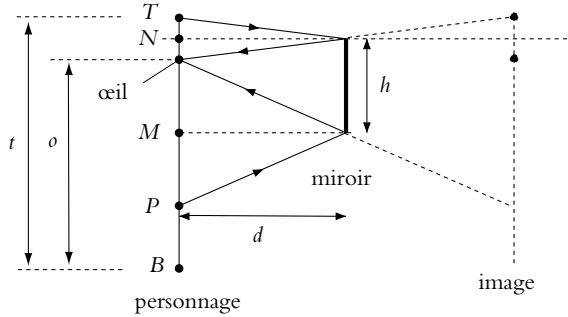
4. La réponse est (a).

La formule de conjugaison s'écrit ici $\frac{1}{SC'} + \frac{1}{SC} = \frac{2}{SC}$ dont on déduit $\overline{SC'} = \overline{SC}$.

5. La réponse est (d).

En effet, on a $\gamma = -\frac{SC'}{SC} = -1$.

- A.4** 1. a) La figure représente le personnage (trait vertical BT) et son image dans le miroir. On a tracé les rayons extrêmes (passant par les bords du miroir) arrivant dans l'œil. Tous les rayons provenant des points situés entre T et P sont donc vus dans le miroir.



- b) Pour savoir si la partie visible du corps dépend de d , on doit calculer la distance TP . On appelle O le point correspondant à la position des yeux. On a donc : $TP = TO + OP = TO + 2 OM$. Or :

$$TO = t - o \quad \text{et} \quad OM = h - NO = h - \frac{TO}{2}$$

soit finalement :

$$TP = 2h$$

La partie du corps visible dans le miroir est donc indépendante de d . Il ne sert à rien de reculer pour voir une plus grande partie de son corps comme on pourrait le penser *a priori*.

- Pour que la personne se voit entièrement dans le miroir, il faut que P et B soient confondus, soit $TP = TB = t$. Avec la relation précédente on en déduit qu'il faut une hauteur h du miroir égale à $\frac{t}{2}$.
- Le fait d'avoir accroché le miroir à mi-distance des yeux et du haut de la tête permet de voir juste le haut de la tête comme le montre la figure.

- A.5** 1. Le miroir domestique est composé d'un dioptre D et d'un miroir plan M_p , ce qui correspond à successivement trois transformations optiques en raison de la réflexion. La suite de conjugaisons est :

$$A \xrightarrow{D} A_1 \xrightarrow{M_p} A_2 \xrightarrow{D} A'$$

2. a) Pour déterminer la position du miroir équivalent M_e , on utilise le fait que M_e se trouve à l'intersection des prolongements des rayons incidents et émergents (figure S10.4).

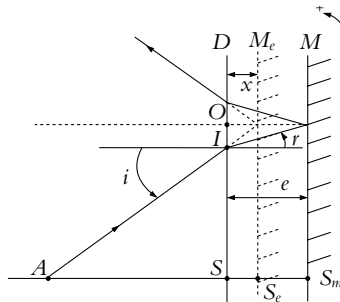


Figure S10.4

- b) On se réfère à la figure (S10.4). On cherche à déterminer $x = KS_e$. On peut écrire (triangle IKS_e) $KI = x \tan i$ et $KI = KS_p \tan r = e \tan r$ (triangle IKS_p). On se place dans l'approximation de Gauss, soit $\tan i \simeq i$ et $\tan r \simeq r$, d'où, $xi = er$. La loi de Descartes en I donne $\sin i = n \sin r$ ou dans l'approximation de Gauss, $i = nr$. En combinant les deux relations, on trouve $x = e/n$.

3. Pour déterminer la position du miroir plan équivalent M_e , on utilise comme propriété que l'image d'un point du miroir est image de lui-même. C'est le point S_e de la figure (S10.4). Si l'on applique la suite de conjugaisons de la question (1) : on cherche l'ensemble des points S_e , images d'eux-mêmes par le système.

$$S_e \xrightarrow{D} A_1 \xrightarrow{M_P} A_2 \xrightarrow{D} S_e$$

On peut utiliser le retour inverse de la lumière pour réécrire la relation précédente :

$$S_e \xrightarrow{D} A_2 \xrightarrow{M_P} A_1 \xrightarrow{D} S_e$$

Par comparaison des deux relations, on déduit que $A_1 = A_2$, donc A_1 est image de lui-même par M_P : il s'agit d'un point de M_P et plus particulièrement S_P . On en déduit que S_e est l'antécédent de S_P par le dioptré D .

On applique la formule de conjugaison du dioptré plan :

$$\frac{1}{\overline{KS_e}} = \frac{n}{\overline{KS_P}} \Rightarrow x = \frac{n}{e}$$

On retrouve bien la même expression que précédemment.

- B.1** 1. Les angles sont représentés sur la figure (S10.5) ; les normales aux miroirs sont en pointillés. La loi de Descartes en O donne $\beta = \varphi + \theta$. L'angle du rayon réfléchi sur M_1 avec l'horizontale est donc $\beta + \theta$. Si l'on se réfère à M_2 , l'angle de ce même rayon avec l'horizontale est $i + \pi/6$. On en déduit : $\beta + \theta = \frac{\pi}{6} + i$ ou $i = 2\theta + \varphi - \pi/6$. Sachant que $\alpha = r - \pi/6$, avec $r = i$ (Descartes), alors $\alpha = 2\theta + \varphi - \pi/3$ l'on se réfère au miroir M_2

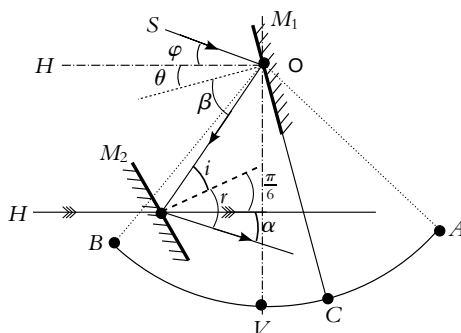


Figure S10.5

2. a) Pour $\alpha = 0$, la relation précédente donne $\theta = \pi/6 - \varphi/2$
 b) Le point V étant le milieu de l'arc AB , l'angle VOC est égal à l'angle θ , donc puisque l'angle AOB vaut $\pi/3$, l'angle AOC est égal à $\pi/6 - \theta$ soit $\varphi/2$.

- B.2** 1. Par définition, le foyer objet a son image à l'infini par le télescope : $F \xrightarrow{M_1} A_i \xrightarrow{M_2} \infty$; A_i est donc le foyer objet de M_2 . On applique la relation de conjugaison avec origine en S_1 :

$$\frac{1}{\overline{S_1 F_2}} + \frac{1}{\overline{S_1 F}} = \frac{2}{R_1}$$

Or la relation de Chasles donne :

$$\overline{S_1 F_2} = \overline{S_1 S_2} + \overline{S_2 F_2} = -D + \frac{R_2}{2}$$

D'où :

$$\overline{S_1 F} = \frac{R_1 \left(-D + \frac{R_2}{2} \right)}{-2D + R_2 - R_1}$$

De manière similaire, pour F' , on a : $\infty \xrightarrow{M_1} A_i \xrightarrow{M_2} F'$. Dans ce cas, A_i est le foyer image de M_1 . On applique la relation de conjugaison avec origine en S_2 :

$$\frac{1}{\overline{S_2 F'_1}} + \frac{1}{\overline{S_2 F'}} = \frac{2}{R_2}$$

Or la relation de Chasles donne :

$$\overline{S_2 F'_1} = \overline{S_2 S_1} + \overline{S_1 F'_1} = D + \frac{R_1}{2}$$

D'où :

$$\overline{S_2 F'} = \frac{R_2 \left(D + \frac{R_1}{2} \right)}{2D + R_1 - R_2}$$

2. a) Le système est afocal si l'image de l'infini est l'infini ; on en déduit la relation suivante :

$$\infty \xrightarrow{M_1} F'_1 = F_2 \xrightarrow{M_2} \infty$$

donc les foyers des deux miroirs doivent être confondus.

- b) La relation de Chasles donne :

$$D = \overline{S_2 S_1} = \overline{S_2 F_2} + \overline{F_2 F'_1} + \overline{F'_1 S_1} = \overline{S_2 F_2} + \overline{F'_1 S_1} = \frac{R_2}{2} - \frac{R_1}{2} = 1,5 \text{ m}$$

- c) Sur la figure (S10.6), on a l'impression que le miroir M_2 occulte le rayon mais dans la réalité, ce miroir est beaucoup plus petit que M_1 . On ne peut respecter l'échelle dans la direction perpendiculaire à l'axe optique sur le schéma, de manière à ce qu'il soit lisible. On a représenté l'image intermédiaire $A_i B_i$.

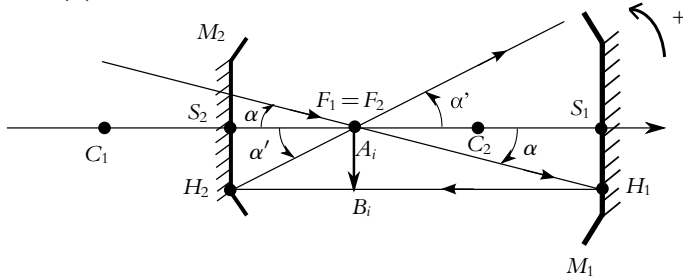


Figure S10.6

Dans l'approximation de Gauss, les angles sont assimilables à leurs tangente. Les relations dans les triangles donnent :

- dans le triangle $S_1 H_1 F_1$: $\tan \alpha \simeq \alpha = -\frac{\overline{S_1 H_1}}{\overline{S_1 F_1}} = -\frac{\overline{A_i B_i}}{\overline{S_1 F_1}}$
- dans le triangle $S_2 H_2 F_2$: $\tan \alpha' \simeq \alpha' = -\frac{\overline{S_2 H_2}}{\overline{S_2 F_2}} = -\frac{\overline{A_i B_i}}{\overline{S_2 F_2}}$

Les signes moins sont dus au fait que $\alpha < 0$ et $\alpha' > 0$, tandis que toutes les grandeurs algébriques sont négatives à l'exception de $\overline{S_2 F_2}$. On en déduit le grossissement :

$$G = \frac{\alpha'}{\alpha} = \frac{\overline{S_1 F_1}}{\overline{S_2 F_2}} = \frac{R_1}{2} \frac{2}{R_2} = -2$$

3. a) L'astre observé et supposé à l'infini, donc l'image finale est au foyer F' . On utilise la relation établie à la question (1). Soit :

$$\overline{S_2 F'} = \frac{R_2 \left(D + \frac{R_1}{2} \right)}{2D + R_1 - R_2} = \frac{-R_1 \left(-2R_1 + \frac{R_1}{2} \right)}{-4R_1 + R_1 + R_1} = -\frac{3}{4} R_1 = \frac{3}{4} R_2$$

L'image intermédiaire A_iB_i est dans le plan focal image de M_1 donc on peut utiliser la relation dans le triangle $F_1S_1H_1$ de la question (2) :

$$\alpha = -\frac{\overline{S_1H_1}}{\overline{S_1F_1}} = -\frac{\overline{A_iB_i}}{\overline{S_1F_1}}$$

Attention, on a raisonné ici sur la demi image de l'astre, puisqu'il est vu symétriquement de part et d'autre de l'axe, sous l'angle $|2\alpha|$. Ensuite, on utilise la formule de grandissement avec origine en S_2 entre l'image intermédiaire A_iB_i et l'image finale $A'B'$:

$$\frac{\overline{A'B'}}{\overline{A_iB_i}} = -\frac{\overline{S_2A'}}{\overline{S_2A_i}} = -\frac{\overline{S_2F'}}{\overline{S_2F_1}}$$

En combinant ces deux relations et en utilisant l'expression de $\overline{S_2F_1}$ de la question (1), il vient :

$$\overline{A'B'} = \alpha \frac{\overline{S_2F'} \overline{S_1F_1}}{\overline{S_2F_1}} = \alpha \frac{-\frac{3}{4}R_1 \frac{R_1}{2}}{-2R_1 + \frac{R_1}{2}} = \alpha \frac{R_1}{2}$$

L'image totale de l'astre a donc une taille deux fois plus grande soit $|\alpha R_1|$

- b) Attention pour l'application numérique, il faut exprimer α en radian puisqu'on a utilisé $\tan \alpha \simeq \alpha$. Soit $2^\circ = 2 \cdot 10^{-2}$ rad. On trouve alors : $\overline{S_2F'} = 1,5$ m et une image de 7 cm.
- c) Sur la figure (S10.7), on trace un rayon incident d'angle α qui passe par F_1 et est donc réfléchi par M_1 parallèlement à l'axe, puis le rayon obtenu est réfléchi par M_2 en passant par F_2 .

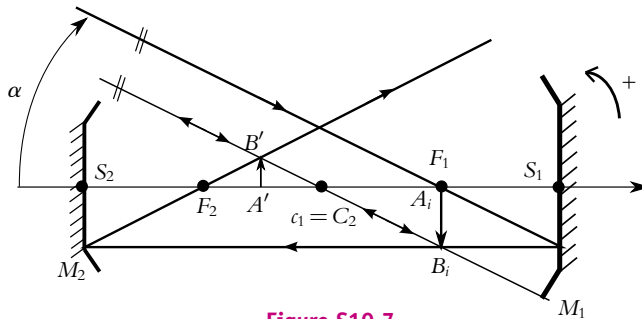


Figure S10.7

L'image intermédiaire est dans le plan de F_1 . Pour obtenir l'image finale, on trace un rayon issu de B_i et passant par C_2 ; ce rayon est réfléchi sur lui-même par M_2 . L'image finale B' est à l'intersection des deux rayons réfléchis par M_2 .

On vérifie bien sur la figure que l'image finale est telle que $\overline{S_2F_2} = \frac{3}{4}R_2$

- B.3** 1. a) On cherche l'antécédent A par le miroir d'une image réelle A' , l'œil de l'observateur. Pour un miroir plan, image et objet sont symétriques par rapport au plan du miroir : l'objet A est donc virtuel. Les points A et A' sont sur l'axe puisqu'un rayon perpendiculaire au miroir est réfléchi sur lui-même.

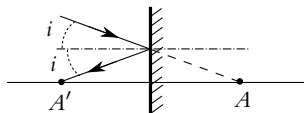


Figure S10.8

On construit de plus (figure S10.8) un rayon réfléchi oblique qui coupe l'axe en A' ; le rayon incident est obtenu par la loi de Descartes de la réflexion et son prolongement coupe l'axe en A .

- b) Pour que l'observateur puisse voir une image par le miroir, il faut que les rayons réfléchis entrent dans son œil, donc interceptent l'axe en A' . D'après la question précédente, tout rayon passant par A' correspond à un incident dont le prolongement passe par A .

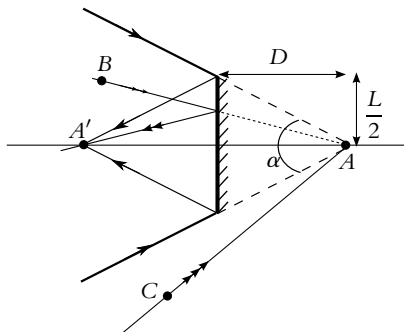


Figure S10.9

Sur la figure (S10.9), on a tracé les rayons limites s'appuyant sur les bords du miroir : ces deux rayons forment un cône de sommet A . D'un point B intérieur au cône, il sera possible de trouver un incident frappant le miroir et dont le prolongement passe par A , ce qui est impossible pour un point C extérieur au cône (le segment AC n'intercepte pas le miroir). L'espace embrassé par le cône à gauche du miroir constitue donc le champ du miroir.

- c) L'angle caractérisant le champ du miroir est l'angle α . D'après la figure (S10.9), on a $\tan \frac{\alpha}{2} = \frac{L}{2D}$.
- d) Application numérique : $\alpha = 22,6^\circ$
2. a) Sur la construction (figure S10.10), l'échelle n'est pas respectée. L'axe optique étant axe de symétrie (système centré), le point A est sur l'axe. Il suffit donc de tracer un autre rayon (oblique). On utilise la propriété que deux rayons réfléchis dont les prolongements se croise en un foyer secondaire F_s correspondent à des rayons incidents parallèles.

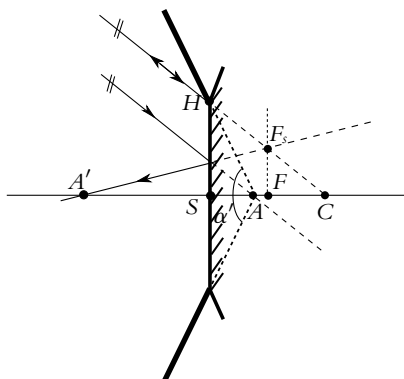


Figure S10.10

Pour déterminer la position de A , on applique la loi de conjugaison avec origine au sommet :

$$\frac{1}{\overline{SA'}} + \frac{1}{\overline{SA}} = \frac{2}{\overline{SC}} \Rightarrow \overline{SA} = \frac{RD}{2D + R} \quad (1)$$

avec $\overline{SA'} = -D$ et $\overline{SC} = R$.

- b) Comme pour un miroir plan, le champ est l'espace embrassé par le cône s'appuyant sur les bords du miroir et de sommet A (traits gras sur la figure (S10.10)).

- c) L'angle α' est obtenu par $\tan \frac{\alpha'}{2} = \frac{SH}{SA}$ avec $SH = \frac{L}{2}$.
- d) Application numérique : $\alpha' = 61,9^\circ$
3. a) Le champ, dans le cas du miroir convexe est considérablement augmenté par rapport au rétroviseur plan. Ceci est dû au fait que le point A est beaucoup plus proche du miroir dans ce cas.
- b) L'objet VW a une image symétrique par rapport au miroir plan (le tracé des rayons a déjà été effectué pour A et A'). En ce qui concerne le miroir convexe, on utilise deux rayons issus de W , l'un parallèle à l'axe, dont le réfléchi a son prolongement passant par F , l'autre dont le prolongement passe par C et qui n'est pas dévié. Les constructions sont réalisées sur les figures (S10.11) et (S10.12).

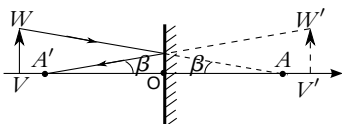


Figure S10.11

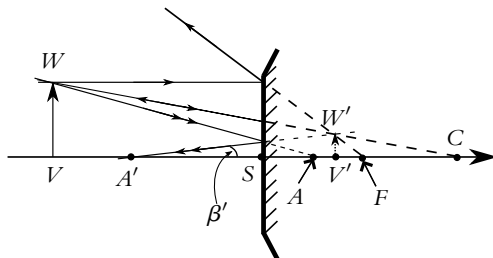


Figure S10.12

On a vu dans les questions précédentes qu'un rayon arrivant dans l'œil après réflexion, correspondait à un rayon incident dont le prolongement passe par A . Les rayons correspondants et les angles β et β' sous lesquels sont vues les images sont représentés sur les constructions. Pour le miroir plan :

$$\tan \beta = \frac{V'W'}{A'W'} = \frac{T}{D + D'} \Rightarrow \beta = 5,4^\circ$$

sachant que objet et image ont la même taille.

Pour le miroir convexe, il faut calculer la position de l'image et sa taille. En ce qui concerne la position, on utilise le résultat de l'équation (1) en remplaçant D par D' , ce qui

$$\text{donne : } \overline{SV'} = \frac{RD'}{2D' + R} = 0,24 \text{ m.}$$

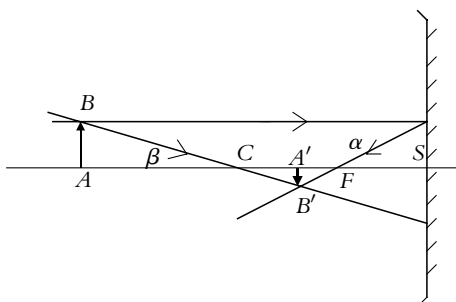
Pour la taille de l'image, en raisonnant en valeur absolue : $\frac{V'W'}{VW} = \frac{SV'}{SV}$, soit

$$V'W' = \frac{SV'}{D'} SW = 2,4 \text{ cm. L'angle } \beta' \text{ est égal à :}$$

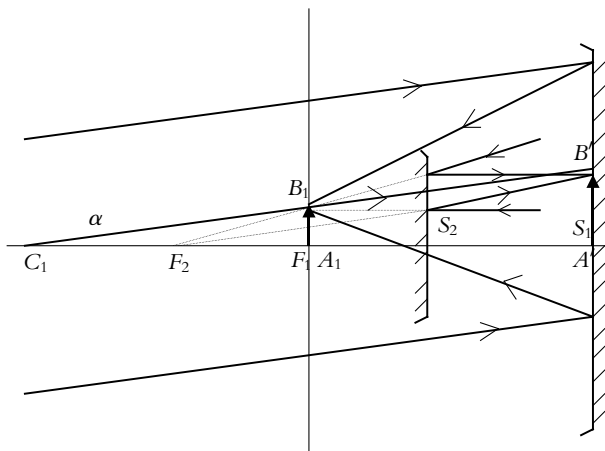
$$\beta' = \arctan \frac{V'W'}{A'V'} = \arctan \frac{V'W'}{D + SV'} = 1,9^\circ$$

L'objet est vu beaucoup plus petit avec le miroir convexe, on a donc l'impression qu'il est plus loin que ce qu'il n'est réellement : il faut en tenir compte lorsqu'on déboîte.

- B.4** 1. Un système est rigoureusement stigmatique pour A et A' si tout rayon passant par A passe par A' .
2. Le stigmatisme est approché si le point est remplacé par un petit volume au voisinage du point. Il est vérifié dans les conditions de Gauss à savoir une faible inclinaison des rayons par rapport à l'axe optique, de faibles angles d'incidence et des rayons arrivant à proximité de l'axe. On note que deux conditions suffisent, la troisième en est une conséquence.
3. Un système est aplanétique si un objet perpendiculaire à l'axe optique donne une image perpendiculaire à l'axe optique.
4. Un rayon passant par le centre C n'est pas dévié car il est dirigé suivant la normale. Un rayon passant par le sommet S ressort symétriquement par rapport à l'axe optique d'après lois de Snell-Descartes. On obtient la figure suivante :



5. On exprime $\tan \alpha$ et $\tan \beta$ par les relations dans un triangle rectangle soit $\tan \alpha = -\frac{\overline{AB}}{\overline{SA}} = \frac{\overline{A'B'}}{\overline{SA'}}$ et $\tan \beta = \frac{\overline{AB}}{\overline{CA}} = \frac{\overline{A'B'}}{\overline{CA'}}$. On en déduit $-\overline{SA'} \overline{CA} = \overline{SA} \overline{CA'}$. En introduisant le point S par la relation de Chasles sur les mesures algébriques, en développant et en divisant par $\overline{SA} \overline{SA'} \overline{SC}$, on obtient $\frac{1}{\overline{SA'}} + \frac{1}{\overline{SA}} = \frac{2}{\overline{SC}}$.
6. Le foyer objet F est le conjugué d'une image à l'infini et le foyer image F' celui d'un objet à l'infini. On en déduit à partir de la relation de conjugaison précédente que F et F' sont confondus au milieu de $[SC]$.
7. On peut considérer que AB est à l'infini donc A_1B_1 est situé dans le plan focal de M_1 passant par F_1 . On obtient B_1 à l'aide d'un rayon incliné d'un angle α passant par le centre C_1 car ce rayon n'est pas dévié. Les autres rayons convergent en B_1 . On construit ensuite $A'B'$ image de A_1B_1 par le miroir M_2 .



8. La relation de conjugaison s'écrit :

$$\frac{1}{\overline{S_2S_1} + \overline{S_1A'}} + \frac{1}{\overline{S_2S_1} + \overline{S_1A_1}} = -\frac{1}{f_2} = \frac{1}{D} + \frac{1}{D - f_1}$$

On en déduit l'équation du second degré en D :

$$D^2 + (2f_2 - f_1)D - f_1f_2 = 0$$

dont la seule solution positive est $D = \frac{-2f_2 + f_1 + \sqrt{(2f_2 - f_1)^2 + 4f_1f_2}}{2}$.

9. Si $f_1 \gg f_2$, on néglige f_2 devant f_1 et on obtient $D \simeq f_1$.

10. On exprime $\tan \alpha$ et on en déduit $\overline{A_1 B_1} = f_1 \tan \alpha$.

11. Par la formule du grandissement, on a $\frac{\overline{A' B'}}{\overline{AB}} = -\frac{\overline{S_2 A'}}{\overline{S_2 A_1}}$ et on en déduit

$$\overline{A' B'} = -f_1 \tan \alpha \frac{-2f_2 + f_1 + \sqrt{(2f_2 - f_1)^2 + 4f_1 f_2}}{-2f_2 - f_1 + \sqrt{(2f_2 - f_1)^2 + 4f_1 f_2}}$$

12. Si $f_1 \gg f_2$, on effectue un développement limité en $\frac{f_2}{f_1}$ et on obtient $\overline{A' B'} = \frac{f_1^2}{f_2} \tan \alpha$.

13. Les applications numériques donnent $\overline{A_1 B_1} = 0,4 \text{ mm}$ et $\overline{A' B'} = 8,02 \text{ mm}$ (on obtiendrait $8,00 \text{ mm}$ par l'expression approchée).

Chapitre 11

A.1 Dans tout l'exercice, on note O le centre optique de la lentille, A l'objet et A' l'image et f' la distance focale image de la lentille.

Si l'image est virtuelle, cela signifie que $\overline{OA'} < 0$. Puisque dans l'exercice on impose une condition sur $\overline{OA}$, on a intérêt à utiliser la relation de conjugaison de Descartes (avec origine en O) :

$$\frac{1}{\overline{OA'}} - \frac{1}{\overline{OA}} = \frac{1}{f'} \quad \text{et donc} \quad \frac{1}{\overline{OA'}} < 0 \Rightarrow \frac{1}{f'} + \frac{1}{\overline{OA}} < 0 \quad \text{soit} \quad \frac{1}{f'} < \frac{1}{\overline{AO}}$$

1. Pour la lentille convergente, comme f' est positive, on en déduit que $\overline{AO}$ est aussi positive donc A est réelle et que $\overline{AO} < f'$ donc A est entre F et O : c'est la loupe.

2. Pour la lentille divergente, f' est négative. Il y a donc deux possibilités :

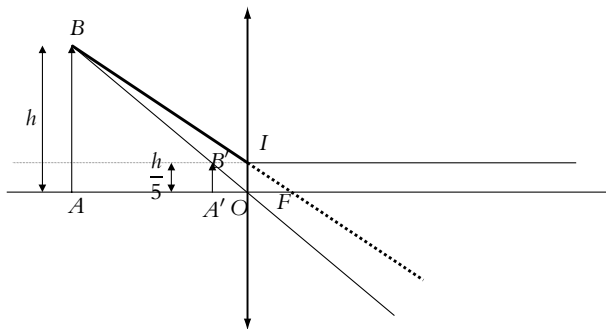
- $\overline{AO} > 0$ et l'inégalité est toujours vérifiée. Tout objet réel donne une image virtuelle.
- $\overline{AO} < 0$ et en faisant attention aux signes $\overline{OA} > -f'$: A est virtuelle et au-delà de F (on rappelle que pour une lentille divergente, le foyer objet est après la lentille.).

A.2 Par le calcul, on utilise la formule de conjugaison avec origine au centre $\frac{1}{\overline{OA'}} - \frac{1}{\overline{OA}} = \frac{1}{f'}$

et l'expression du grandissement avec origine au centre $\gamma = \frac{\overline{OA'}}{\overline{OA}}$. L'image se trouve en A' tel que

$\overline{OA'} = \gamma \overline{OA} = -12 \text{ cm}$ et la distance focale est $f' = \frac{\gamma \overline{OA}}{1 - \gamma} = -15 \text{ cm}$. La lentille est donc divergente.

Pour la construction, on trace le rayon sortant parallèlement à l'axe optique à une distance de l'axe 5 fois plus petite que celle de l'objet. Ce rayon est censé être passé par le foyer objet. Puis on trace le rayon passant par le centre qui n'est pas dévié. L'image B' est à l'intersection de ces deux rayons ou de leur prolongation. On obtient le foyer objet comme l'intersection de l'axe optique et du rayon passant par B et I (intersection du rayon sortant parallèle à l'axe optique déjà tracé et de la lentille).



- A.3** 1. Puisque la lentille est divergente, on remarque que l'objet A est confondu avec le foyer image de la lentille. Pour trouver l'image A' , on applique par exemple la relation de conjugaison de Newton : $\overline{F'A'} \cdot \overline{FA} = -f^2$. Avec $\overline{FA} = 2f' = -2f$, on trouve : $\overline{F'A'} = f/2$, donc A' est à mi-distance de F' et O .
2. Pour retrouver la formule de grandissement, on applique la relation de Thalès dans le triangle OAB de la figure (S11.1), ce qui donne : $\frac{\overline{A'B'}}{\overline{AB}} = \frac{\overline{OA'}}{\overline{OA}} = \frac{1}{2}$ soit $\overline{A'B'} = 0,5 \text{ mm}$.

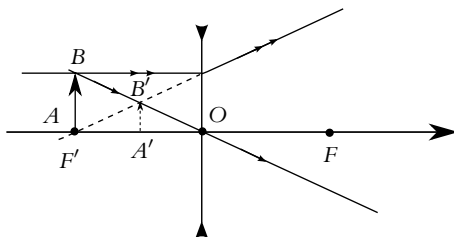


Figure S11.1

- A.4** 1. La réponse est (b) car un faisceau provenant de l'infini repart à l'infini, ce qui définit un système afocal.
2. La réponse est (a). Un point à l'infini donne une image en F'_1 par L_1 qui donne une image en F_3 par L_2 pour obtenir par L_3 une image à l'infini. On applique la relation de conjugaison de la lentille L_2 soit $\frac{1}{\overline{O_2F_3}} + \frac{1}{\overline{O_2F'_1}} = \frac{1}{f_2}$ avec $\overline{O_2F'_1} = f_1 - \Delta$ et $\overline{O_2F_3} = \delta - f_3$.
3. La réponse est (d). On a $\frac{r_3}{f_3} = \frac{r_2}{F_3O_2}$ donc $r_2 = \frac{r_3}{f_3}(f_3 - \delta)$ et de même $\frac{r_2}{F'_1O_2} = \frac{-r_1}{f_1}$ donc $r_2 = -\frac{r_1}{f_1}(f_1 - \Delta)$. On en déduit $\frac{r_3}{r_1} = \frac{f_3}{f_1} \frac{f_1 - \Delta}{-f_3 + \delta}$.
4. La réponse est (a). On utilise les relations des deux questions précédentes et on en déduit $\Delta = f_1 + f_2 \left(1 - \frac{r_3 f_1}{r_1 f_3}\right)$.
5. La réponse est (c). On reporte l'expression de Δ obtenue à la question précédente dans celle du rapport $\frac{r_3}{r_1}$ et on obtient $\delta = f_3 + f_2 \left(1 - \frac{r_1 f_3}{r_3 f_1}\right)$.

- A.5** 1. L'objet étant à l'infini, l'image est dans le plan focal image de L , c'est donc dans ce plan que l'on doit placer la pellicule.

2. a) On applique la relation de conjugaison de Descartes avec $\overline{OA} = -d_A$, soit :

$$\frac{1}{\overline{OA'}} - \frac{1}{-d_A} = \frac{1}{f'} \Rightarrow \overline{OA'} = \frac{d_A f'}{d_A - f'}$$

- b) Sur la figure (S11.2) on trace deux rayons incidents s'appuyant sur le bord de la lentille. On trace ensuite le rayon passant par le centre de la lentille et parallèle à l'un des rayons incidents.

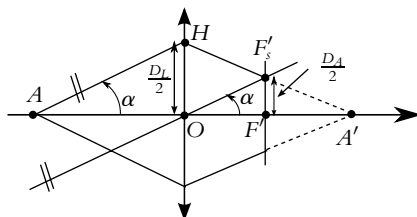


Figure S11.2

Ces deux rayons donnent des rayons émergents qui se croisent en F'_S dans le plan focal image de L . L'image A' se trouve donc au-delà de la pellicule qui est dans le plan focal de L . On peut alors écrire, dans le triangle OAH :

$$\tan \alpha = \frac{D_L}{2 \overline{OA}} = \frac{D_L}{2d_A}$$

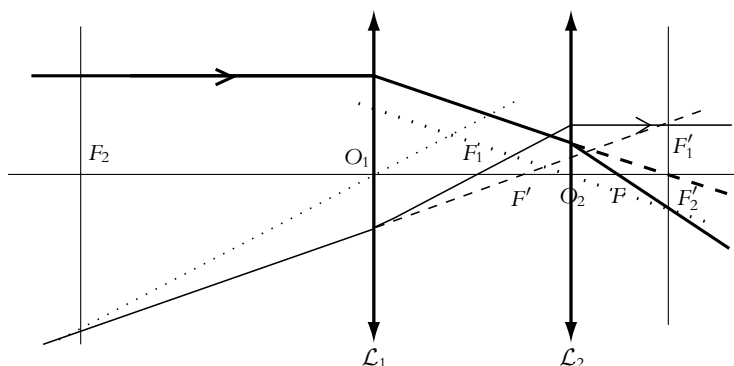
et dans le triangle $OF'F'_S$:

$$\tan \alpha = \frac{D_A}{2f'}$$

d'où la relation : $D_A = D_L \frac{f'}{d_A}$.

- c) Pour que l'image soit nette, il faut que $D_A \leq \epsilon$, soit $d_A \geq f' \frac{D_L}{\epsilon}$, ou $d_A \geq 3$ m.

- A.6** 1. F' est l'image d'un point de l'axe optique à l'infini. Un rayon arrivant de l'infini passe par F'_1 et est dévié par L_2 . Son trajet est déterminé par un foyer secondaire obtenu comme l'intersection du plan focal image de L_2 et le rayon parallèle à celui entre L_1 et L_2 et passant par le centre de L_2 .



2. L'infini a pour image F'_1 par $\mathcal{L}_1$ qui a pour image F' par $\mathcal{L}_2$. On applique la relation de conjugaison de Newton pour $\mathcal{L}_2$ et on utilise la relation de Chasles pour obtenir $\overline{F'_2 F'} = \frac{f_2^2}{-f'_1 + e + f'_2} = -\frac{a}{2}$.
3. F est l'objet d'un point de l'axe optique à l'infini et on applique le même principe qu'à la première question pour déterminer sa position. Cf. première question pour la construction.

4. L'infini a pour antécédent F_2 par $\mathcal{L}_2$ qui a pour antécédent F par $\mathcal{L}_1$. On applique la relation de conjugaison de Newton pour $\mathcal{L}_1$ ainsi que la relation de Chasles pour obtenir $\overline{F_1 F} = \frac{-f_1'^2}{-f_1' + e + f_2'} = \frac{9a}{2}$.

A.7 1. La formule de conjugaison s'écrit $\frac{1}{\overline{O_1 A_1}} - \frac{1}{\overline{O_1 A_0}} = \frac{1}{f_1'}$ et l'expression du grandissement est $\gamma = \frac{\overline{A_1 B_1}}{\overline{A_0 B_0}} = \frac{\overline{O_1 A_1}}{\overline{O_1 A_0}}$. On en déduit $\overline{O_1 A_0} = \frac{(1 - \gamma)f_1'}{\gamma} = -20$ cm.

2. La position de l'image est donnée par $\overline{O_1 A_1} = \gamma \overline{O_1 A_0} = -10$ cm.

3. Pour avoir une image sur l'écran de l'objet initial, il faut que $\mathcal{L}_2$ donne une image de $A_1 B_1$ sur l'écran.

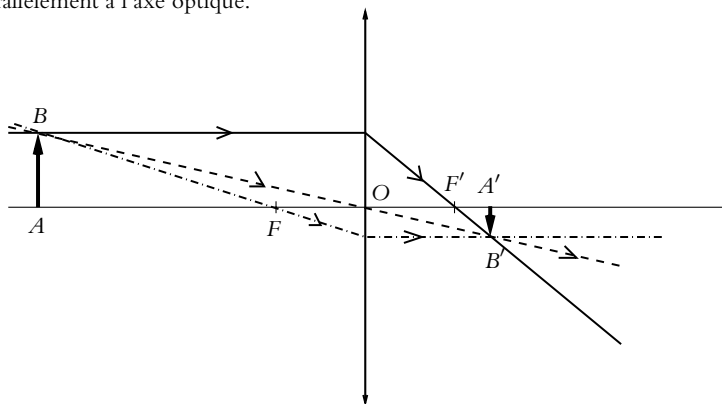
En appliquant la formule de conjugaison, on obtient $\overline{O_1 O_2} = \frac{\overline{O_2 E_2'} f_2'}{\overline{O_2 E} - f_2'} + \overline{O_1 A_1} = 70$ cm.

4. Pour réaliser l'opération souhaitée, le système doit être afocal donc $F_1' = F_2$. Par conséquent, on en déduit que $\overline{O_1 O_2} = 20$ cm.

5. On détermine l'expression de $\tan \alpha$ en notant α l'angle du rayon entre les deux lentilles issu de F_1' et l'axe optique soit $\tan \alpha = \frac{d}{-f_2'} = \frac{D}{f_1'}$. On en déduit $\frac{D}{d} = -\frac{f_2'}{f_1'} = 2$.

B.1 1. L'infini est conjugué avec le foyer image F' . Pour l'autre borne en A à 1,40 m, on a une image en A' tel que $\frac{1}{\overline{OA'}} - \frac{1}{\overline{OA}} = \frac{1}{f'}$ soit $\overline{OA'} = \frac{\overline{OA} f'}{f' + \overline{OA}} = 93,3$ mm. Les images sont donc entre 93,3 et 100 mm, c'est dans cette marge qu'on doit placer la pellicule.

2. On trace les rayons suivants : un rayon passant par le centre de la lentille n'est pas dévié, un rayon arrivant parallèlement à l'axe passe par le foyer image F' et un rayon passant par le foyer objet F ressort parallèlement à l'axe optique.



3. Le grandissement s'écrit $\gamma = \frac{\overline{A'B'}}{\overline{AB}} = \frac{\overline{OF}}{\overline{AF}} = \frac{\overline{A'F'}}{\overline{OF'}}$.

On en déduit $\overline{A'B'} = \overline{AB} \frac{\overline{OF}}{\overline{AF}} = -\frac{f' \overline{AB}}{\overline{AO} - f'}$. Les applications numériques donnent une taille de $-2,0$ mm pour la tour et $-7,7$ cm pour l'objet d'un mètre à 1,40 m de l'objectif.

4. a) On cherche les points B_1 et B_2 dont les images par $\mathcal{L}_1$ sont respectivement l'infini et un point à 1,40 m de l'objectif. On utilise la relation de conjugaison $\frac{1}{\overline{O_1 B_i}} - \frac{1}{\overline{O_1 A_i}} = \frac{1}{f_1'}$ dont on tire

$\overline{O_1 B_i} = \frac{\overline{O_1 A_i} f_1'}{f_1' - \overline{O_1 A_i}}$. Les applications numériques donnent $\overline{O_1 B_2} = -267$ mm pour A à 1,40 m ($\overline{O_1 A} = \overline{O_1 O} + \overline{OA} = 50 - 1400$ soit -1350 mm) et $\overline{O_1 B_2} = -333$ mm pour A à l'infini.

- b) L'image d'un objet AB à 1,00 m est par L_1 $\overline{O_1A_1} = \frac{f'_1 \overline{O_1A}}{f'_1 + \overline{O_1A}} = -3,03$ m puis par L $\overline{OA'} = \frac{f' \overline{OA_1}}{f' + \overline{OA_1}} = 0,103$ m. La formule du grandissement $\gamma = \frac{\overline{OA'}}{\overline{OA}}$ appliquée successivement aux deux lentilles donne un grandissement de 10 pour L_1 et de $-3,37 \cdot 10^{-2}$ pour L soit une taille $A'B' = -0,337$ m.
- c) La taille est 4,72 fois plus grande : la deuxième lentille agrandit les objets proches.
5. a) On procède comme à la question 4.a et on obtient qu'on peut photographier de 22,1 m à l'infini pour une même course de l'objectif.
- b) On cherche les positions des images par L_3 puis L_2 puis L et on applique à chaque fois la formule du grandissement avec origine au centre de la lentille. On obtient $\gamma_3 = -3,33 \cdot 10^{-5}$, $\gamma_2 = -7,50$ donc $\gamma = \gamma_1 \gamma_2 \gamma_3 = -1,33 \cdot 10^{-4}$.
- c) On a un grandissement très faible donc on peut photographier un panorama plus vaste.

B.2 1. On a $\overline{O_1A'} = d$ et $\overline{O_1A} = \overline{O_1A'} + \overline{A'A} = d - D$ qu'on reporte dans la relation de conjugaison avec origine au centre $\frac{1}{\overline{O_1A'}} - \frac{1}{\overline{O_1A}} = \frac{1}{f'_1}$. En réduisant au même dénominateur, on obtient l'équation $d^2 - Dd + f'_1 D = 0$.

2. La résolution de cette équation donne $d = \frac{D \pm \sqrt{D^2 - 4f'_1 D}}{2}$. Les applications numériques donnent 49,9 m et 5,01 cm. La première valeur est impossible car elle est supérieure à 100 mm ; la seconde est acceptable.

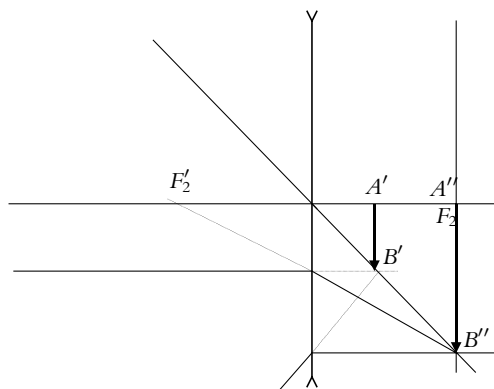
3. L'expression du grandissement s'écrit

$$\gamma = \frac{\overline{A'B'}}{\overline{AB}} = \frac{\overline{OF}}{\overline{AF}} = \frac{-f'_1}{D - d - f'_1} = -1,0 \cdot 10^{-3}.$$

4. On en déduit $h' = 2,0$ mm.

5. L'application de la relation de conjugaison à la deuxième lentille permet d'écrire $\frac{1}{\overline{O_2A''}} - \frac{1}{\overline{O_2A'}} = \frac{1}{f'_2}$ soit en utilisant la relation de Chasles $\overline{O_2A'} = \overline{O_2A''} + \overline{A''A'}$, on en déduit $\overline{A''A'} = \frac{d_2^2}{f'_2 - d_2} = -2,0 \cdot 10^{-2}$ m.

6. On trace le rayon passant par le centre qui n'est pas dévié, le rayon arrivant parallèlement à l'axe optique qui passe par le foyer image F' et le rayon passant par le foyer objet qui ressort parallèlement à l'axe optique.



7. Le calcul du grandissement donne $\gamma = \frac{\overline{A'B'}}{\overline{AB}} = \frac{\overline{F'A'}}{\overline{OF'}} = \frac{-f'_2 + d_2}{f'_2} = -2$.
8. On effectue la même résolution qu'aux deux premières questions en remplaçant D par $D' = D + \overline{A''A'} = 49,998 \text{ m}$ et d' par $d'' = d' + \overline{A''A'} = d' - 2 \cdot 10^{-3}$. On obtient $d'' = \frac{D \pm \sqrt{D^2 - 4f'_1 D}}{2}$. Les deux valeurs obtenues sont 49,95 m et 5,01 cm. La première est impossible du fait de l'encombrement et la seconde convient donc $d' = 5,21 \text{ cm}$.
9. On procède de même pour le grandissement $\gamma_1 = \frac{-f'_1}{D' + d'' - f'_1} = -10^{-3}$ et $\gamma = \gamma_1 \gamma_2 = 2 \cdot 10^{-3}$ donc $h'' = 4,0 \text{ mm}$.
10. Le négatif et l'objet restent à la même place mais le grandissement double d'où le nom de doubleur de focale donné au dispositif.

B.3 1. a) L'objet est réel et l'image se formant sur la rétine (qui joue le rôle d'écran) est donc réelle elle-aussi. La seule possibilité d'objet réel donnant une image réelle correspond au cas d'une lentille convergente (ici le cristallin de l'œil), avec l'objet avant F et l'image après F' et dans ce cas l'image est renversée.

- b) On connaît les distances de l'objet ($\overline{OA} = -D$) et de l'image au centre optique, donc on utilise la formule de grandissement avec origine en O :

$$\gamma = \frac{\overline{A'B'}}{\overline{AB}} = \frac{\overline{OA'}}{\overline{OA}} = \frac{d}{D} = 11 \cdot 10^{-3}$$

On en déduit la taille de l'image : $\overline{A'B'} = \gamma \overline{AB} = 1,1 \text{ mm}$.

- c) Pour calculer la vergence du système, on applique la relation de Descartes :

$$V = \frac{1}{\overline{OA'}} - \frac{1}{\overline{OA}} = \frac{1}{d} - \frac{1}{-D} = 92 \delta$$

2. a) On ne peut se servir des résultats des questions précédentes. Le fait que la position de la rétine (position de l'image) soit inchangée alors que la position de l'objet ait changée implique que la vergence V de la lentille ait évolué. Le raisonnement est le même qu'à la question (a) : image réelle renversée puisqu'elle se forme sur un écran (la rétine) et que l'objet est réel.

- b) On applique de nouveau la loi de Descartes :

$$V = \frac{1}{\overline{OA'}} - \frac{1}{\overline{OA}} = \frac{1}{11 \cdot 10^{-3}} - \frac{1}{-25 \cdot 10^{-2}} = 491 \delta$$

On trouve que le cristallin s'est contracté (la vergence a augmenté) pour voir plus près. La taille de l'image est obtenue par l'expression du grandissement :

$$\overline{A'B'} = \gamma \overline{AB} = \frac{\overline{OA'}}{\overline{OA}} \overline{AB} = \frac{0,11}{-0,25} 0,1 = -4,4 \text{ mm}$$

3. a) On calcule tout d'abord la vergence V du cristallin de l'œil myope. L'image d'un objet à l'infini se forme dans le plan focal image de la lentille donc à une distance f' de O soit $f' = 11 - 0,5 = 10,5 \text{ mm}$ soit $V = 95,2 \delta$. On appelle L_D le verre de lunette et R le point d'intersection de l'axe optique avec la rétine. Pour que le myope voit net l'objet à l'infini, on doit avoir la suite de conjugaison suivante : $\infty \xrightarrow{L_D} F'_D \xrightarrow{\text{cristallin}} R$. Il s'agit donc de trouver la position de l'antécédent de R par le cristallin. Avec $\overline{OR} = d$, la relation de Descartes donne :

$$V = \frac{1}{d} - \frac{1}{\overline{OF'_D}} \Rightarrow \overline{OF'_D} = \frac{d}{1 - Vd} = -23,3 \text{ cm}$$

On remarque que F'_D est à 23,3 cm avant O donc à 21,3 cm avant L_D ce qui correspond bien à une lentille divergente de distance focale image $f'_D = -21,3 \text{ cm}$, donc de vergence $V' = 1/f'_D = -4,7 \delta$.

- b) Dans ce cas, la vergence du cristallin a changé par rapport à la question précédente, puisque l'œil n'observe plus à l'infini. On sait cependant que l'image finale A' est sur la rétine ($\overline{OA'} = d$). On calcule la position de l'image intermédiaire A_i par la formule Descartes appliquée à L_D :

$$\frac{1}{\overline{OA_i}} - \frac{1}{\overline{OA'}} = V' \Rightarrow \overline{OA_i} = \frac{\overline{OA'}}{1 + \overline{OA'} V'} = -17,5 \text{ cm}$$

On en déduit $\overline{OA_i} = -19,5 \text{ cm}$.

Le grandissement de l'ensemble est le produit des grandissements :

$$\gamma = \frac{\overline{O'A_i}}{\overline{OA}} \frac{\overline{OA'}}{\overline{OA_i}} = -0,01$$

- B.4** 1. a) Un système est dit afocal s'il conjugue l'infini avec l'infini, c'est-à-dire qu'un objet à l'infini a une image à l'infini.

L'image (intermédiaire) de l'infini par L_1 est dans le plan focal image de L_1 . Par L_2 , cette image intermédiaire a une image à l'infini, donc elle doit être dans le plan focal objet de L_2 , d'où la série des conjugaisons :

$$\infty \xrightarrow{L_1} F'_1 = F_2 \xrightarrow{L_2} \infty$$

- b) Schéma de la lunette :

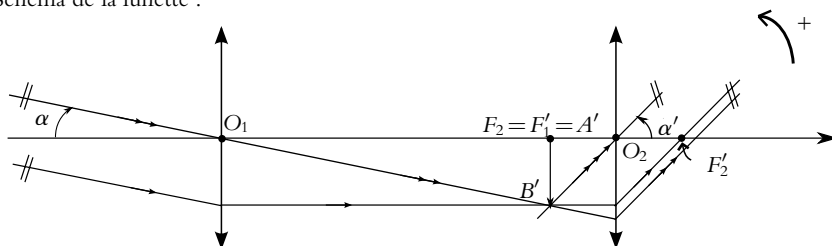


Figure S11.3

- c) L'oculaire (L_2) est utile pour observer à l'œil (qui observe à l'infini pour ne pas accommoder). Si l'on souhaite prendre une photo, on place la pellicule dans le plan focal image de l'objectif L_1 .
2. a) Sur la figure (S11.3), on remarque que pour un faisceau lumineux incident provenant du haut, le faisceau émergent semble provenir du bas : l'image est donc renversée, ce qui est confirmé par le sens différent des angles α et α' .
- b) Dans les triangles $O_1A'B'$ et $O_2A'B'$ (on rappelle que $A' = F_2 = F'_1$), on écrit :

$$\tan \alpha = \frac{\overline{A'B'}}{\overline{O_1A'}} = \frac{\overline{A'B'}}{f'_1} \quad \text{et} \quad \tan \alpha' = -\frac{\overline{A'B'}}{\overline{A'O_2}} = -\frac{\overline{A'B'}}{f'_2}$$

On note que $\alpha < 0$, $\alpha' > 0$ et $\overline{A'B'} < 0$. Dans l'approximation de Gauss, $\tan \alpha \simeq \alpha$ et $\tan \alpha' \simeq \alpha'$, soit

$$G = \frac{\alpha'}{\alpha} = -\frac{f'_1}{f'_2} = -5$$

- c) Le verre a un indice de réfraction n qui dépend de la longueur d'onde λ en suivant la loi de Cauchy : $n = a + \frac{b}{\lambda^2}$, où a et b sont des constantes propres au matériau.

La vergence d'une lentille dépend de l'indice du verre (elle est proportionnelle à $n - 1$, si la lentille est utilisée dans l'air) en raison des réfractions sur les faces ; elle est donc plus grande pour le bleu que pour le rouge et donc la lentille est plus convergente pour le bleu que pour le rouge. La conséquence dans le cas d'une lumière blanche est un décalage entre les images de différentes couleurs et l'apparition d'irisations colorées.

Dans le cas d'un miroir, la réflexion (métallique) est indépendante de la longueur d'onde dans le visible.

3. a) Par la lunette, on a toujours, un objet à l'infini dont l'image par l'objectif L_1 est dans le plan focal image de L_1 et une image finale à l'infini, dont l'antécédent par l'oculaire L_2 est dans le plan focal objet de L_2 . Il faut donc que la lentille intermédiaire L_3 conjugue les points F'_1 et F_2 .
- b) En appliquant la formule du grandissement avec origine au centre optique de L_3 pour F'_1 et F_2 , on peut écrire :

$$\gamma_3 = \frac{\overline{O_3 F_2}}{\overline{O_3 F'_1}} \Rightarrow \overline{O_3 F_2} = \gamma_3 \overline{O_3 F'_1}$$

On reporte dans la relation de conjugaison avec origine au centre pour L_3 :

$$\frac{1}{\overline{O_3 F_2}} - \frac{1}{\overline{O_3 F'_1}} = \frac{1}{f'_3} \quad \text{soit} \quad \frac{1}{\overline{O_3 F'_1}} \left(\frac{1}{\gamma_3} - 1 \right) = \frac{1}{f'_3}$$

d'où :

$$\overline{O_3 F'_1} = f'_3 \left(\frac{1 - \gamma_3}{\gamma_3} \right)$$

- c) Schéma avec les trois lentilles :

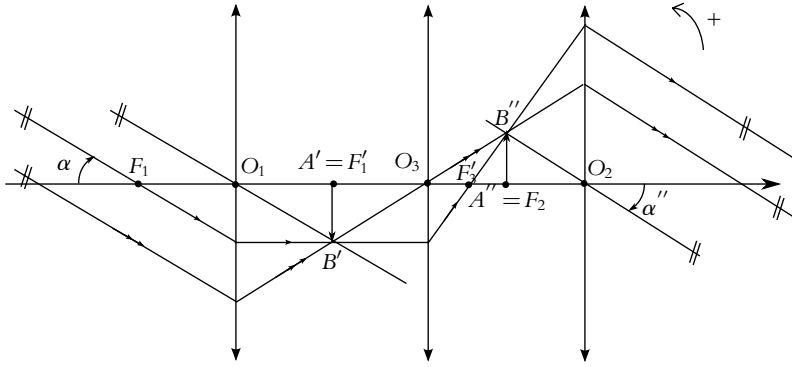


Figure S11.4

- d) En utilisant les notations de la figure (S11.4), on obtient :

- Pour L_1 , on a : $\tan \alpha \simeq \alpha = \frac{\overline{A'B'}}{f'_1}$;
- pour L_3 , on a : $\gamma_3 = \frac{\overline{A''B''}}{\overline{A'B'}}$;
- pour L_2 , on a : $\tan \alpha'' \simeq \alpha'' = -\frac{\overline{A''B''}}{f'_2}$.

On vérifie les signes des expressions : les deux angles sont négatifs ainsi que $\overline{A'B'}$ et γ_3 , les autres grandeurs étant positives. En combinant ces équations, on obtient :

$$G' = \frac{\alpha''}{\alpha} = \gamma_3 \left(-\frac{f'_1}{f'_2} \right)$$

Puisque γ_3 est négatif (figure S11.4), G' est bien positif et l'image est à l'endroit. En ce qui concerne la valeur absolue de G' cela dépend de celle de γ_3 . Si $|\gamma_3| > 1$ alors $G' > |G|$.

- B.5** 1. A l'œil nu, un objet à une distance d_m et de taille ℓ , est vu sous l'angle α_m , tel que :

$$\tan \alpha_m = \frac{\ell}{d_m}$$

avec $\tan \alpha_m \simeq \alpha_m$ dans l'approximation de Gauss : $\alpha_m = \frac{\ell}{d_m}$.

2. a) La lentille est utilisée en loupe, donc l'objet est placé entre le foyer objet F et le centre optique O . On note $x = \overline{OA}$ et $y = \overline{OA'}$, A' étant l'image de A . La relation de conjugaison de Descartes donne :

$$\frac{1}{y} - \frac{1}{x} = \frac{1}{f'} \Rightarrow y = \frac{xf'}{f' + x}$$

On note O_b la position de l'œil de l'observateur. Pour que l'image soit vue nette il faut, $O_b A' \geq d_m$, soit en valeur algébrique : $-a + y \leq -d_m$. Avant de chercher l'intervalle des positions x de l'image, il faut connaître les variations de y en fonction de x . On calcule donc la dérivée :

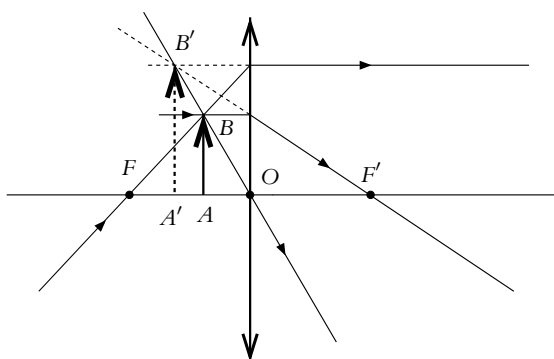
$$\frac{dy}{dx} = \frac{1}{(f' + x)^2} (f'(f' + x) - xf') = \frac{(f')^2}{(f' + x)^2} > 0$$

donc y est une fonction croissante de x : à l'objet le plus proche correspond l'image la plus proche, et à l'objet le plus lointain ($A = F$) correspond l'image à l'infini. L'objet le plus proche, vu net par l'œil est tel que $y = a - d_m$, soit :

$$\frac{1}{a - d_m} - \frac{1}{x} = \frac{1}{f'} \Rightarrow x = \frac{f'(a - d_m)}{f' - a + d_m}$$

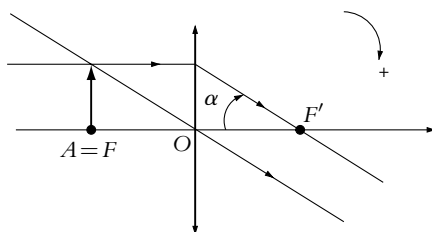
La zone possible pour les objets donnant une vision nette est donc :

$$-f' \leq x \leq \frac{f'(a - d_m)}{f' - a + d_m}$$



- b) L'œil n'accomode pas si l'image A' est à l'infini, donc l'objet doit être dans le plan focal de la lentille : $A = F$. D'après la figure ci-contre :

$$\tan \alpha \simeq \alpha = \frac{\ell}{f'}$$



On déduit de ce qui précède le grossissement commercial de la loupe :

$$G = \frac{\alpha}{\alpha_m} = \frac{\ell}{f'} \frac{d_m}{\ell} \Rightarrow G = \frac{d_m}{f'} = 5$$

- B.6** 1. L'oculaire joue le rôle de la loupe étudiée précédemment (problème 5.) et le réticule le rôle de l'objet A , mais $d = -x$:

$$\left| \frac{f'(a - d_m)}{f' - a + d_m} \right| \leq d \leq f'$$

Pour qu'il n'y ait pas d'accommodation, l'image du réticule doit être à l'infini, donc le réticule doit être placé au foyer objet F_2 de L_2 . Les constructions correspondent aux deux précédentes.

2. a) D'après l'énoncé, l'image de A à travers L_1 est O , alors :

$$\frac{1}{\overline{O_1O}} - \frac{1}{\overline{O_1A}} = \frac{1}{f'_1} \Rightarrow (\overline{O_1A} - \overline{O_1O})f'_1 = \overline{O_1A} \cdot \overline{O_1O}$$

Soit en faisant apparaître $x = \overline{OA}$:

$$\overline{O_1A}^2 - x\overline{O_1A} - xf'_1 = 0$$

Les solutions de ce trinôme en $\overline{O_1A}$ doivent être réelles, donc le discriminant doit être positif :

$$\Delta = x^2 + 4xf'_1 \geq 0 \Rightarrow x(x + 4f'_1) \geq 0$$

or $x < 0$ donc $x + 4f'_1 \leq 0$ et $x \leq -4f'_1$.

- b) A partir de l'équation de conjugaison de Descartes, on peut établir un trinôme en D comme on l'a fait précédemment pour $\overline{O_1A}$:

$$D^2 + xD - xf'_1 = 0$$

On garde pour D la solution positive (attention $x < 0$) :

$$D = -\frac{x}{2} - \frac{1}{2}\sqrt{x^2 + 4xf'_1}$$

- c) On a démontré dans l'étude de la loupe (problème 5.) que $D = \overline{O_1A_1}$ est une fonction croissante de $\overline{O_1A}$. Dans le cas présent, $\overline{O_1A} < 0$ et $\overline{O_1A_1} > 0$, donc la valeur minimale de D est obtenue pour $\overline{O_1A} \rightarrow -\infty$ soit $D = f'_1$, et la valeur maximale pour l'objet le plus proche de L_1 . En calculant la dérivée D par rapport à x , on peut montrer que D est une fonction croissante de x , donc D sera maximale pour $x = -4f'_1$, soit $D_{\max} = 2f'_1$. d'où :

$$f'_1 \leq D \leq 2f'_1$$

3. a) Quand l'observateur change, seul le réglage de l'oculaire est à modifier. On part d'un viseur réglé pour un œil emmétrope à l'infini. Dans ce cas, le réticule est dans le plan focal objet de l'oculaire donc l'image finale est à l'infini. On souhaite augmenter la distance d , de manière à ce que l'image finale soit au *Punctum Remotum* de l'observateur. En appliquant la relation de conjugaison de Descartes :

$$\frac{1}{\overline{O_2A'}} - \frac{1}{\overline{O_2A}} = \frac{1}{f'_2}$$

or $\overline{O_2A'} = a - \delta$, ce qui donne :

$$d = \overline{OO_2} = \frac{(\delta - a)f'_2}{\delta - a + f'_2}$$

- b) Il suffit pour chaque œil de calculer la position du *Punctum Remotum*. Lorsqu'on place des verres correcteurs contre l'œil, l'image d'un objet à l'infini par le verre correcteur est dans le plan focal image de ce verre. Comme le verre est accolé à l'œil, la distance du P.R. à l'œil est égale à la distance focale.
- pour l'œil myope avec correction de -8δ , $\delta = 12,5$ cm et $d = 1,72$ cm ;
 - pour l'œil hypermétrope avec correction de $+8 \delta$, $\delta = -12,5$ cm et $d = 2,38$ cm.

Chapitre 12

A.1 Un objet à l'infini donne une image au foyer par le miroir. L'œil verra cette image si elle est située entre le PP et le PR soit entre 15 et 40 cm : $15 \leq \overline{OF} \leq 40$ cm.

La distance entre le foyer et le sommet vaut la moitié du rayon du miroir $\overline{FS} = \frac{R}{2}$. Le miroir est à 70 cm au maximum donc $0 \leq \overline{OS} = \overline{OF} + \overline{FS} \leq 70$ cm. On en déduit $-\overline{OF} \leq \overline{FS} \leq 70 - \overline{OF}$ avec $15 \leq \overline{OF} \leq 40$ soit $-40 \leq -\overline{OF} \leq -15$. Par conséquent, $-40 \leq \overline{FS} \leq 55$ cm et $-80 \leq R \leq 110$ cm.

- A.2** 1. Un objet A a pour image A_1 par L_1 et A' par l'ensemble des deux lentilles. Si les lentilles sont accolées, on peut confondre leurs deux centres qu'on note O dans la suite. Les relations de conjugaison donnent $\frac{1}{OA_1} - \frac{1}{OA} = \frac{1}{f'_1}$ et $\frac{1}{OA'} - \frac{1}{OA_1} = \frac{1}{f'_2}$ soit en sommant les deux $\frac{1}{OA'} - \frac{1}{OA} = \frac{1}{f'_1} + \frac{1}{f'_2}$ qui est la formule de conjugaison d'une lentille de vergence $V = \frac{1}{f'} = V_1 + V_2 = \frac{1}{f'_1} + \frac{1}{f'_2}$.

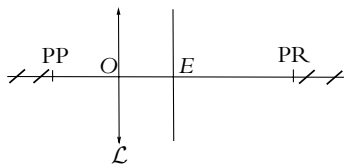
2. Par application de la relation proposée, on a

$$V_F - V_C = (n_F - 1)A - (n_C - 1)A = (n_F - n_C)A$$

Or la formule du pouvoir dispersif permet d'écrire $n_F - n_C = \frac{n_D - 1}{\nu}$. On en déduit $V_F - V_C = \frac{n_D - 1}{\nu}A = \frac{V_D}{\nu}$. En remplaçant la vergence par la distance focale, on a $\frac{1}{f'_F} - \frac{1}{f'_C} = \frac{1}{\nu f'_D}$. En supposant que $f'_D \simeq f'_C \simeq f'_F$, on peut écrire $\frac{f'_C - f'_F}{f'_C f'_F} \simeq \frac{f'_C - f'_F}{f'^2_D} = \frac{1}{\nu f'_D}$ soit $f'_C - f'_F = \frac{f'_D}{\nu} = 6,0 \text{ mm}$. Cette valeur est très faible devant f'_D donc l'approximation est justifiée.

3. Les conséquences sont une irisation de l'image mais une moins bonne netteté.
4. On souhaite $V_F = V_C$ soit $V_{1F} + V_{2F} = V_{1C} + V_{2C}$ qui s'écrit en explicitant les expressions de la vergence $A_1(n_{1F} - n_{1C}) = A_2(n_{2C} - n_{2F})$ avec $A_i = \frac{V_{iD}}{n_{iD} - 1}$. Après simplification, on obtient $\frac{V_{1D}}{\nu_1} + \frac{V_{2D}}{\nu_2} = 0$.
5. On n'a rendu le doublet de lentilles achromatique uniquement pour les raies C et F . Les autres ne le sont pas mais on a réduit l'écart. On pourra considérer le système comme approximativement achromatique.
6. On veut $V_{1D} + V_{2D} = 2,0$ et $\frac{V_{1D}}{\nu_1} + \frac{V_{2D}}{\nu_2} = 0$. La résolution de ce système donne $V_{1D} = \frac{2\nu_2}{\nu_2 - \nu_1} = 4,0$ et $V_{2D} = \frac{2\nu_1}{\nu_1 - \nu_2} = -2,0$ soit $f'_1 = -50 \text{ cm}$ et $f'_2 = 25 \text{ cm}$.

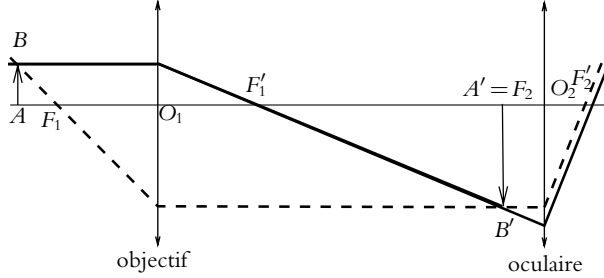
- B.1** 1. a)



- b) Les objets sont réels de l'infini à PP et virtuels de PR à l'infini.
- c) On modélise ainsi le défaut d'hypermétropie.
- d) Au repos, le point A est en PR et A' en E soit en appliquant la formule de conjugaison avec origine au centre, on a $f' = \frac{dD_m}{D_m - d} > 0$. La lentille est donc convergente.
- e) L'accommodation est maximale pour A en PP et A' en E soit par le même principe $f' = \frac{d_m d}{d_m - d}$.
- f) On accole les deux lentilles donc on confond les deux centres. La combinaison des formules de conjugaison permet d'obtenir $f' = \frac{f'_1 f'_2}{f'_1 + f'_2}$ avec un centre en $O = O_1 = O_2$.
- g) L'œil normal est au repos si A est à l'infini et A' en E. On en déduit $f' = d = \frac{f'_c f'_o}{f'_c + f'_o}$ avec $f'_o = \frac{dD_m}{D_m - d}$ soit $f'_c = D_m$. La lentille est donc convergente.

L'œil normal accommode au maximum si A est au PP et A' en E . On en déduit $f' = \frac{da}{a-d} = \frac{f'_c f'_o}{f'_c + f'_o}$ avec $f'_o = \frac{dd_m}{d_m - d}$ soit $f'_c = \frac{ad_m}{a - d_m}$. La lentille est donc convergente.

2. a)



b) Avec l'oculaire, on a $\alpha' = \frac{\overline{A'B'}}{\overline{O_2 F_2'}} = \frac{\overline{A'B'}}{f'_2}$. Sans oculaire, on a $\alpha = \frac{\overline{A'B'}}{\delta}$. On en déduit que $f'_2 = \frac{\delta}{g}$.

c) A a pour image F_2 par $\mathcal{L}_1$ puis l'infini par $\mathcal{L}_2$. On en déduit que $G = \frac{\overline{F_1 F_2}}{\overline{F_1 O_1}}$ et que $f'_1 = \frac{-\Delta}{G}$.

d) Le grandissement s'écrit $G = \frac{\overline{O_1 F_2}}{\overline{O_1 A}} = \frac{f'_1 + \Delta}{G}$ dont on tire $\overline{O_1 A} = -\frac{\Delta}{G^2} + \frac{\Delta}{G}$.

e) Le cas de l'œil accommodant à l'infini a été traité à la question précédente.

Le cas de l'œil accommodant à δ est tel que A_1 a pour image A_2 par $\mathcal{L}_1$ puis A_3 par $\mathcal{L}_2$ avec $\overline{F_2' A_3} = \delta$. En tenant compte de la formule de conjugaison avec origine aux foyers, on en déduit

$$\overline{O_1 A_1} = -f'_1 - \frac{f_1'^2}{\Delta - \frac{f_2'^2}{\delta}}.$$

Par conséquent, la latitude de mise au point s'écrit

$$\overline{O_1 A_1} - \overline{O_1 A} = -f'_1 - \frac{f_1'^2}{\Delta - \frac{f_2'^2}{\delta}} + \frac{\Delta}{G^2} - \frac{\Delta}{G}.$$

f) On a $\alpha' = \frac{\overline{A'B'}}{f'_2}$ et $\alpha = \frac{\overline{AB}}{\delta}$. On en déduit $g_c = G \frac{\delta}{f'_2} = Gg$.

g) On fait l'approximation que $u \simeq 45^\circ$. On a donc de forts angles d'incidence et les conditions de Gauss ne sont pas respectées : on observe les défauts de sphéricité. On a $D = 2\overline{O_1 A} \tan u$.

h) Le point O_1 a pour image C par l'oculaire soit $F_2' C = \frac{f_2'^2}{\Delta + f'_1}$. La taille du cercle oculaire D' est

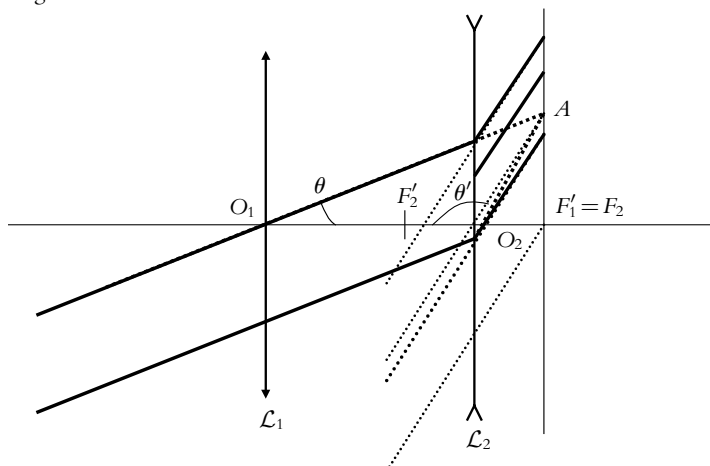
$$\text{telle que } \frac{D'}{D} = \frac{\overline{F_2' O_2}}{\overline{F_2' O_1}} \text{ soit } D' = \frac{f_2' D}{\Delta + f'_1}$$

Par conséquent, le cercle oculaire est le plus petit cercle où arrivent tous les rayons. On a donc un maximum de lumière.

B.2 1. La lentille $\mathcal{L}_1$ est convergente de distance focale $f'_1 = \frac{1}{V_1} = 20$ cm, la lentille $\mathcal{L}_2$ est divergente de distance focale $f'_2 = \frac{1}{V_2} = -5,0$ cm.

2. Pour avoir un système afocal, il ne faut pas de foyer autrement l'infini est conjugué de l'infini. Cela impose d'avoir $F'_1 = F_2$ et la distance entre les lentilles est $d = f'_1 + f'_2$ soit 15 cm.

3. Un rayon passant par le centre de la lentille n'est pas dévié et s'il vient de l'infini, il converge dans le plan focal image.



4. On a $\theta' = \frac{\overline{F_2 A}}{\overline{O_2 F_2}}$ et $\theta = \frac{\overline{F_1' A}}{\overline{O_1 F_1'}}$. On en déduit $\gamma = \frac{\theta'}{\theta} = -\frac{f_1'}{f_2'} = 4,0$.
5. L'angle de vision sans lunette est $\frac{D}{L}$ soit $2,5 \cdot 10^{-4}$ pour Copernic et $6,25 \cdot 10^{-4}$ pour Clavius. Par conséquent, Copernic n'est pas visible à l'œil tandis que Clavius l'est. Le grossissement valant 4,0, les deux seront visibles avec lunette.
6. Le grandissement d'une lentille s'exprime par $\frac{\overline{A'B'}}{\overline{AB}} = -\frac{\overline{F'A'}}{\overline{OF'}} = \frac{\overline{OF}}{\overline{FA}}$ soit ici en appliquant cette relation aux deux lentilles, on en déduit $G = -\frac{f_2'}{f_1'}$.

Deuxième période

Partie IV – Électrocinétique II

Chapitre 15

- A.1** 1. En série, on a $\underline{Z}' = R' + jL'\omega$ et en parallèle,

$$\underline{Z} = \frac{RjL\omega}{R + jL\omega} = \frac{RL^2\omega^2}{R^2 + L^2\omega^2} + j\frac{R^2L\omega}{R^2 + L^2\omega^2}.$$

En identifiant les parties réelles et imaginaires, on a l'équivalence des deux dipôles pour :

$$R' = \frac{RL^2\omega^2}{R^2 + L^2\omega^2} \quad \text{et} \quad L' = \frac{R^2L}{R^2 + L^2\omega^2}$$

On peut noter que l'équivalence n'est valable que pour une fréquence précise. On verra dans le chapitre sur l'instrumentation électrique qu'il s'agit du modèle bande étroite d'une bobine.

2. Le rapport $\frac{L'}{R'}$ vaut, d'après la question précédente :

$$\frac{L'}{R'} = \frac{R^2L(R^2 + L^2\omega^2)}{(R^2 + L^2\omega^2)RL^2\omega^2} = \frac{R}{L\omega^2}$$

On aura $\frac{L'}{R'} = \frac{L}{R}$ si

$$\frac{R}{L\omega^2} = \frac{L}{R} \quad \text{soit pour} \quad \omega = \frac{R}{L}$$

A.2 1. On transforme l'association de l'inductance L et de la source idéale de tension en un modèle de Norton d'impédance $jL\omega$ et de courant de court-circuit $\frac{\underline{E}}{jL\omega}$. On associe ensuite la capacité C

et l'inductance L en parallèle, on obtient une impédance $\frac{jL\omega \frac{1}{jC\omega}}{jL\omega + \frac{1}{jC\omega}} = \frac{jL\omega}{1 - LC\omega^2}$. On transforme

alors les deux modèles de Norton en modèles de Thévenin : l'un d'impédance $\frac{jL\omega}{1 - LC\omega^2}$ et de force

électromotrice $\frac{\underline{E}}{1 - LC\omega^2}$, l'autre d'impédance R' et de force électromotrice $R'\underline{I}_0$. On remarque que les forces électromotrices sont en opposition. Il suffit alors d'appliquer une loi des mailles pour obtenir

$$\underline{i} = \frac{\underline{E} - R'(1 - LC\omega^2)\underline{I}_0}{(R + R')(1 - LC\omega^2) + jL\omega}.$$

2. On choisit la masse sur la ligne du bas et on appelle A le nœud entre R et R' , B le nœud entre R' et C . On calcule le potentiel en A en appliquant le théorème de Millman soit

$$\underline{v}_A = \frac{\frac{\underline{v}_B}{R'} - \underline{I}_0}{\frac{1}{R} + \frac{1}{R'}} = \frac{R\underline{v}_B - RR'\underline{I}_0}{R + R'}. \text{ De même, pour le potentiel en } B, \text{ on a}$$

$$\underline{v}_B = \frac{\underline{I}_0 + \frac{\underline{v}_A}{R'} + \frac{\underline{E}}{jL\omega}}{\frac{1}{R'} + jC\omega + \frac{1}{jL\omega}} = \frac{R'\underline{E} + jR'L\omega\underline{I}_0 + jL\omega\underline{v}_A}{R'(1 - LC\omega^2) + jL\omega}.$$

En reportant l'expression de $\underline{v}_B$ dans celle de $\underline{v}_A$, on obtient après simplification

$$\underline{v}_A = \frac{R(\underline{E} - R'(1 - LC\omega^2)\underline{I}_0)}{(R + R')(1 - LC\omega^2) + jL\omega} \text{ et } \underline{i} = \frac{\underline{v}_A}{R}.$$

A.3 On calcule l'impédance équivalente à l'ensemble L , R_2 et C soit

$$\underline{Z} = \frac{\frac{1}{jC\omega}(R_2 + jL\omega)}{\frac{1}{jC\omega} + R_2 + jL\omega} = \frac{R_2 + jL\omega}{1 + jR_2C\omega - LC\omega^2}$$

On en déduit l'intensité $\underline{I}$ par une loi des mailles $\underline{I} = \frac{\underline{E}}{R_1 + \underline{Z}}$. Intensité et tension sont en phase si l'argument de $\underline{Z}$ est nul soit $R_2^2C = L(1 - LC\omega^2)$.

B.1 1. a) On a donc une impédance (C_p) en parallèle sur deux impédances en série (L et C_s). On calcule l'admittance de l'ensemble :

$$\underline{Y} = jC_p\omega + \frac{1}{jL\omega + \frac{1}{jC_s\omega}} = jC_p\omega + \frac{jC_s\omega}{1 - LC_s\omega^2}$$

En réduisant au même dénominateur puis en factorisant, on obtient :

$$\underline{Y} = j(C_p + C_s)\omega \frac{1 - \frac{LC_pC_s\omega^2}{C_p + C_s}}{1 - LC_s\omega^2}$$

En inversant, on peut écrire $\underline{Z}$ sous la forme de l'énoncé :

$$\underline{Z}_{AB} = \frac{1}{\underline{Y}} = \left(-\frac{j}{\alpha\omega} \right) \frac{1 - \frac{\omega^2}{\omega_r^2}}{1 - \frac{\omega^2}{\omega_a^2}}$$

avec $\alpha = C_p + C_s$, $\omega_r = 1/\sqrt{LC_s}$ et $\omega_a = 1/\sqrt{LC_e}$ où l'on a posé $C_e = C_p C_s / (C_p + C_s)$.
Comme $C_e < C_s$, alors $\omega_a > \omega_r$.

- b) La relation entre pulsation et fréquence étant $\omega = 2\pi f$, l'application numérique donne :
 $f_r = 796$ Hz et $f_a = 800$ kHz.
2. a) L'impédance $\underline{Z}_{AB}$ est un imaginaire pur, donc la phase φ est $\pm\pi/2$. Si $\varphi = \pi/2$ le comportement est inductif (comme une bobine), sinon il est capacitif (comme un condensateur). Pour cette étude, il suffit d'étudier le signe du rapport

$$-\frac{1 - \frac{\omega^2}{\omega_r^2}}{1 - \frac{\omega^2}{\omega_a^2}}$$

On a représenté φ en fonction de ω (figure S15.1) :

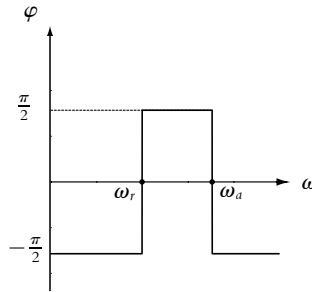


Figure S15.1

Le comportement est donc capacitif pour $\omega \in [0, \omega_r]$ ou $\omega \in [\omega_a, \infty[$ et inductif pour $\omega \in [\omega_r, \omega_a]$.

- b) Sur la figure (S15.2), on a tracé l'allure du module Z_{AB} en fonction de ω .

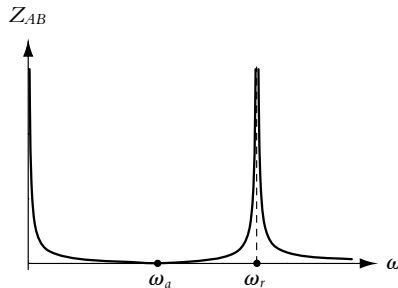


Figure S15.2

- c) Si l'on tient compte de la résistance de la bobine, les valeurs de ω_a et ω_r seront légèrement différentes. Il n'y aura pas de pulsation qui annule totalement $|Z|$ ou qui le rende infini.

B.2 1. L'impédance de L , R_p et C_p en série est : $\underline{Z}_1 = jL\omega + R_p + \frac{1}{jC\omega}$. Cette impédance est en parallèle avec le condensateur C soit :

$$\underline{Y} = jC\omega + \frac{1}{\underline{Z}_1} = jC\omega + \frac{1}{jL\omega + R_p + \frac{1}{jC\omega}} = jC\omega + \frac{R_p - j\left(L\omega - \frac{1}{C_p\omega}\right)}{R_p^2 + \left(L\omega - \frac{1}{C_p\omega}\right)^2}$$

On veut que la partie imaginaire de $\underline{Y}$ soit nulle, ce qui entraîne :

$$C\omega = \frac{\left(L\omega - \frac{1}{C_p\omega}\right)}{R_p^2 + \left(L\omega - \frac{1}{C_p\omega}\right)^2} \Rightarrow C = \frac{\left(L - \frac{1}{C_p\omega^2}\right)}{R_p^2 + \left(L\omega - \frac{1}{C_p\omega}\right)^2}$$

2. On obtient alors :

$$R_1 = \frac{1}{\underline{Y}} = \frac{R_p^2 + \left(L\omega - \frac{1}{C_p\omega}\right)^2}{R_p} = R_p + \frac{1}{R_p} \left(L\omega - \frac{1}{C_p\omega}\right)^2$$

B.3 1. le dipôle D est en série avec C , donc l'impédance équivalente est $\underline{Z}_1 = \underline{Z} + \frac{1}{jC\omega}$. Ce dipôle est lui-même en parallèle sur L , soit un dipôle d'admittance équivalente $\underline{Y}_2 = \underline{Y}_1 + \underline{Y}$ soit

$$\underline{Y}_2 = \frac{1}{jL\omega} + \frac{1}{\underline{Z} + \frac{1}{jC\omega}}$$

Ce dipôle est lui-même en série avec C ce qui donne finalement un impédance :

$$\underline{Z}_e = \frac{1}{jC\omega} + \frac{1}{\underline{Y}_2} = \frac{1}{jC\omega} + \frac{1}{\frac{1}{jL\omega} + \frac{1}{\underline{Z} + \frac{1}{jC\omega}}}$$

2. On souhaite imposer $\underline{Z}_e = \underline{Z}$. On peut alors écrire :

$$\underline{Z}_e = \frac{1}{jC\omega} + \frac{jL\omega \left(\underline{Z} + \frac{1}{jC\omega}\right)}{\underline{Z} + \frac{1}{jC\omega} + jL\omega} = \frac{\underline{Z} + \frac{1}{jC\omega} + jL\omega + jL\omega(1 + \underline{Z}jC\omega)}{1 + \underline{Z}jC\omega - LC\omega^2}$$

La condition $\underline{Z}_e = \underline{Z}$ donne alors :

$$\underline{Z} + \underline{Z}^2 jC\omega - LC\omega^2 \underline{Z} = \underline{Z} + \frac{1}{jC\omega} + jL\omega + jL\omega - \underline{Z}LC\omega^2$$

soit après simplification :

$$\underline{Z}^2 = -\frac{1}{C^2\omega^2} + \frac{2L}{C}$$

Suivant le signe de $\underline{Z}^2$ on a :

$$\text{si } \underline{Z}^2 > 0 \text{ alors } \underline{Z} = \sqrt{-\frac{1}{C^2\omega^2} + \frac{2L}{C}} \text{ et si } \underline{Z}^2 < 0 \text{ alors } \underline{Z} = j\sqrt{\frac{1}{C^2\omega^2} - \frac{2L}{C}}$$

3. a) Aux hautes fréquences, $(C\omega)^2 \gg \frac{C}{2L}$, alors $\underline{Z} \simeq \sqrt{2L/C}$. L'impédance est réelle ce qui revient à considérer que la ligne est équivalente à une résistance, ce qui n'est pas le cas à basse fréquence car $\underline{Z}^2 < 0$ et $\underline{Z}$ est imaginaire pur.

Si l'on se donne un facteur 100 entre les deux termes alors :

$$(C\omega)^2 > 100 \frac{C}{2L} \Rightarrow \omega > 10 \sqrt{\frac{1}{2LC}} \Rightarrow \omega > 50\,000 \text{ rad.s}^{-1} \text{ soit } f > 7,9 \text{ kHz}$$

b) La résistance R est égale à $\sqrt{2L/C}$ soit $R = 100 \, \Omega$.

c) Comme on vient de l'étudier, pour $\omega \rightarrow +\infty$, le circuit se comporte comme une résistance. En revanche, pour $\omega = 0$, alors $Z_e \rightarrow \infty$, donc le circuit est équivalent à un interrupteur ouvert.

B.4 1. On écrit une loi des mailles dans chaque circuit soit

$$L \frac{d^2 q_1}{dt^2} + M \frac{d^2 q_2}{dt^2} + R \frac{dq_1}{dt} + \frac{q_1}{C} = -u \quad \text{et} \quad L \frac{d^2 q_2}{dt^2} + M \frac{d^2 q_1}{dt^2} + R \frac{dq_2}{dt} + \frac{q_2}{C} = 0$$

2. On passe à la notation complexe et on en déduit

$$(1 - LC\omega^2 + jRC\omega) \underline{q_1} - MC\omega^2 \underline{q_2} = -C\underline{u}$$

et

$$-MC\omega^2 \underline{q_1} + (1 - LC\omega^2 + jRC\omega) \underline{q_2} = 0$$

On a donc

$$\underline{q_1} = \frac{-C\underline{u} (1 - LC\omega^2 + jRC\omega)}{(1 - LC\omega^2 + jRC\omega)^2 - M^2 C^2 \omega^2}$$

et

$$\underline{q_2} = \frac{-MC^2 \omega^2 \underline{u}}{(1 - LC\omega^2 + jRC\omega)^2 - M^2 C^2 \omega^2}$$

3. On résout ce système, ce qui permet d'écrire $\underline{i_1} = j\omega \underline{q_1} = \frac{(R + jX) \underline{u}}{(X - jR)^2 - M^2 \omega^2}$ et $\underline{i_2} = j\omega \underline{q_2} = \frac{-jM\omega \underline{u}}{(X - jR)^2 - M^2 \omega^2}$.

4. L'amplitude des courants est donnée par le module des expressions qui viennent d'être obtenues soit

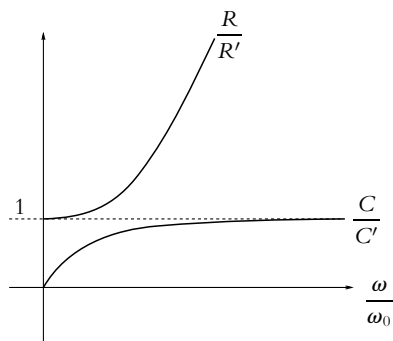
$$I_1 = \frac{\sqrt{X^2 + R^2} U}{\sqrt{4R^2 X^2 + (X^2 - R^2 - M^2 \omega^2)^2}}$$

et

$$I_2 = \frac{M\omega U}{\sqrt{4R^2 X^2 + (X^2 - R^2 - M^2 \omega^2)^2}}$$

B.5 1. Les impédances s'écrivent $\underline{Z} = \frac{1}{jC\omega + \frac{1}{R}} = \frac{R}{1 + jRC\omega}$ et $\underline{Z}' = R' + \frac{1}{jC\omega}$. On obtient en identifiant $R' = \frac{R}{1 + (RC\omega)^2}$ et $C' = C \frac{1 + (RC\omega)^2}{(RC\omega)^2}$.

2. On en déduit $\frac{R}{R'} = 1 + \left(\frac{\omega}{\omega_0}\right)^2$ et $\frac{C}{C'} = \frac{\left(\frac{\omega}{\omega_0}\right)^2}{1 + \left(\frac{\omega}{\omega_0}\right)^2}$. Les allures sont les suivantes :

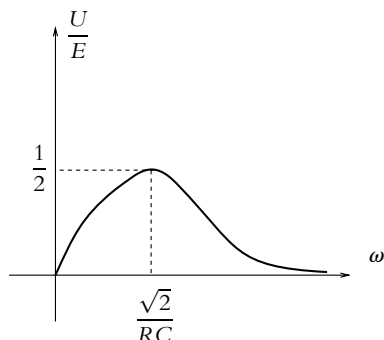


3. On utilise la relation du pont diviseur de tension et on obtient $\frac{u}{\varepsilon} = \frac{jRC\omega}{2 + 2jRC\omega - R^2C^2\omega^2}$ ainsi que $\frac{v}{\varepsilon} = \frac{(1 + jRC\omega)^2}{2 + 2jRC\omega - R^2C^2\omega^2}$.

4. Pour l'amplitude, on prend le module soit $U = \frac{RC\omega}{\sqrt{(2 - R^2C^2\omega^2)^2 + 4R^2C^2\omega^2}} E$.

5. On pose $x = RC\omega$ et on étudie la fonction $f(x) = \frac{x}{\sqrt{4 + x^4}}$ dont la dérivée vaut $f'(x) = \frac{4 - x^4}{(4 + x^4)^{\frac{3}{2}}}$.

La fonction f est croissante de 0 à $\frac{1}{2}$ pour $\omega < \frac{\sqrt{2}}{RC}$ et décroissante ensuite jusqu'à 0.

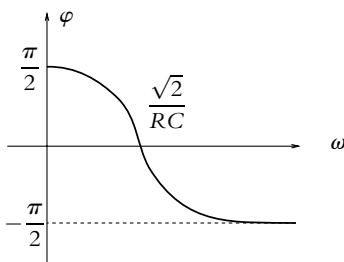


6. On en déduit le déphasage

$$\varphi = \begin{cases} \frac{\pi}{2} - \arctan \frac{2RC\omega}{2 - R^2C^2\omega^2} & \text{si } \omega < \frac{\sqrt{2}}{RC} \\ -\frac{\pi}{2} - \arctan \frac{2RC\omega}{2 - R^2C^2\omega^2} & \text{si } \omega > \frac{\sqrt{2}}{RC} \end{cases}$$

en tenant compte des signes des cosinus et des sinus.

7. On pose $x = RC\omega$ et on étudie $f(x) = \frac{2x}{2 - x^2}$ dont la dérivée vaut $f'(x) = \frac{4 + 2x^2}{(2 - x^2)^2} > 0$. La fonction f est croissante et on en déduit que φ est une fonction décroissante de $+\frac{\pi}{2}$ à $-\frac{\pi}{2}$ en passant par 0 pour $\omega = \frac{\sqrt{2}}{RC}$.

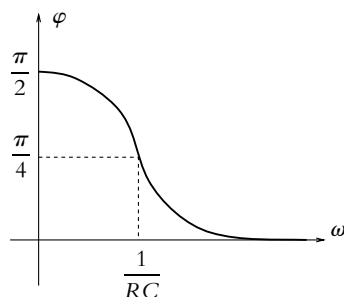


8. L'intensité complexe est $i = \frac{\mathcal{E}}{Z_{\text{tot}}} = \frac{jC\omega(1 + jRC\omega)}{1 - R^2C^2\omega^2 + 3jRC\omega}\mathcal{E}$.

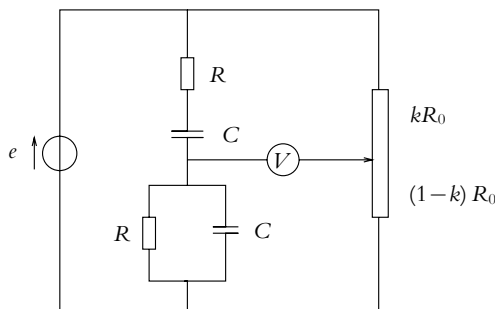
9. Le déphasage s'écrit $\varphi = \psi_1 - \psi_2$ avec $\psi_1 = \pi - \arctan \frac{1}{RC\omega}$ qui est une fonction croissante de $\frac{\pi}{2}$ à π et

$$\psi_2 = \begin{cases} \arctan \frac{3RC\omega}{1 - R^2C^2\omega^2} & \text{si } \omega < \frac{1}{RC} \\ \pi + \arctan \frac{3RC\omega}{1 - R^2C^2\omega^2} & \text{si } \omega > \frac{1}{RC} \end{cases}$$

en tenant compte des signes des cosinus et des sinus. La fonction ψ_2 est croissante de 0 à π en passant par $\frac{\pi}{2}$ en $\omega = \frac{1}{RC}$. Finalement la fonction φ est décroissante de $\frac{\pi}{2}$ à 0 en passant par $\frac{\pi}{4}$ en $\omega = \frac{1}{RC}$.



10. Il n'y a pas de déphasage entre u et e pour $\omega = \omega_0$, ce qui sera aussi le cas aux bornes de $(1 - k)R_0$. C'est la raison pour laquelle on choisit cette pulsation.



11. On veut également que les modules des tensions u et celle aux bornes de $(1 - k)R_0$ soient égales soit $(1 - k)E = \frac{E}{\sqrt{5}}$. On en déduit $k = \frac{\sqrt{5} - 1}{\sqrt{5}}$.

Chapitre 16

A.1 1. L'association en parallèle de R , L et C admet pour impédance

$$\underline{Z} = \frac{1}{\frac{1}{R} + \frac{1}{jL\omega} + jC\omega} = \frac{jRL\omega}{R + jL\omega - RCL\omega^2}$$

On obtient alors l'expression de la tension s en utilisant la relation du pont diviseur de tension soit

$$\underline{s} = \frac{\underline{Z}}{\underline{Z} + r} \underline{e} = \frac{jRL\omega}{rR + jL\omega(r + R) - rRCL\omega^2} \underline{e}.$$

2. On a résonance si l'amplitude passe par un maximum. Or l'amplitude s'obtient en déterminant le module soit $|\underline{s}| = \frac{|\underline{e}|}{\sqrt{\left(1 + \frac{r}{R}\right)^2 + r^2 \left(C\omega - \frac{1}{L\omega}\right)^2}}$. Le dénominateur est toujours plus grand que

$1 + \frac{r}{R}$ et prend cette valeur pour $\omega_0 = \frac{1}{\sqrt{LC}}$. On a donc résonance pour ω_0 .

3. La bande passante s'obtient en résolvant $|\underline{s}| \leq \frac{|\underline{s}|_{\max}}{\sqrt{2}}$ avec $|\underline{s}|_{\max} = \frac{R|\underline{e}|}{r + R}$. La résolution s'effectue comme dans le cours et on obtient $\omega_1 \leq \omega \leq \omega_2$ avec

$$\omega_1 = \frac{-(r + R)L + \sqrt{(r + R)^2 L^2 + 4R^2 r^2 LC}}{2rRLC}$$

et

$$\omega_2 = \frac{(r + R)L + \sqrt{(r + R)^2 L^2 + 4R^2 r^2 LC}}{2rRLC}$$

On en déduit $\Delta\omega = \frac{r + R}{rRC}$.

4. Le facteur de qualité s'obtient par $Q = \frac{\omega_0}{\Delta\omega} = \frac{rR}{r + R} \sqrt{\frac{C}{L}}$.

5. Pour $\omega = \omega_0$, on a $\underline{s} = \frac{R\underline{e}}{r + R}$ et il n'y a pas de déphasage.

6. On obtient la même pulsation de résonance que pour le circuit R, L, C série mais un facteur de qualité variable en fonction de la valeur de r .

A.2 1. On aura un montage du type pont diviseur de tension avec deux impédances $\underline{Z}_1$ et $\underline{Z}_2$. Par conséquent, on a $\underline{s} = \frac{\underline{Z}_2}{\underline{Z}_1 + \underline{Z}_2} \underline{e}$. On veut un gain de 1 à la résonance, ce qui offre deux possibilités : $\underline{Z}_1(\omega_0) = 0$ ou $\underline{Z}_2(\omega_0)$ infinie.

Pour le premier cas, on associe en série L et C soit $\underline{Z}_1 = j\left(L\omega - \frac{1}{C\omega}\right)$ qui s'annule pour $\omega_0 = \frac{1}{\sqrt{LC}}$ soit $C = \frac{1}{L\omega_0^2} = 158 \text{ nF}$. Dans ce cas, on a $\underline{Z}_2 = R$.

Pour le deuxième cas, $\underline{Z}_2$ s'obtient par une association en parallèle de L et C soit $\underline{Z}_2 = \frac{1}{j\left(C\omega - \frac{1}{L\omega}\right)}$ qui tend vers l'infini pour $\omega = \omega_0$. Dans ce cas, $\underline{Z}_1 = R$.

2. Pour étudier la sélectivité, il faut déterminer la bande passante et le facteur de qualité.

Pour le premier cas, par une étude analogue à celle du cours, on trouve $Q = \frac{1}{R} \sqrt{\frac{L}{C}}$.

Pour le second $Q = R \sqrt{\frac{C}{L}}$.

La sélectivité est meilleure pour le second montage car on peut augmenter la valeur de R tandis que diminuer sa valeur est difficile.

3. On a déjà déterminé la valeur $C = 158 \text{ nF}$.

Pour le premier cas, on prend la valeur la plus faible pour R soit $1 \text{ k}\Omega$ pour ne pas être gêné par les résistances internes des sources de tension. On en déduit $Q = 0,50$.

Pour le second cas, on choisit la valeur la plus grande possible pour R soit $100 \text{ k}\Omega$ pour ne pas être gêné par les impédances d'entrée des appareils de mesure. Dans ce cas, $Q = 199$.

4. Le deuxième montage est donc meilleur d'après ce qui précède.

- B.1** 1. a) On note Z_{eq} l'impédance équivalente à R et C en parallèle et $Y_{eq} = \frac{1}{R} + jC\omega$ l'admittance. Le diviseur de tension entre $\underline{u}$ et $\underline{e}$ donne :

$$\underline{u} = \frac{Z_{eq}}{jL\omega + Z_{eq}} \underline{e} = \frac{1}{j\omega Y_{eq} + 1} \underline{e} = \frac{E\sqrt{2}e^{j\omega t}}{1 - LC\omega^2 + \frac{jL\omega}{R}}$$

- b) On transforme le dipôle $(e(t), L)$ en série (modèle de Thévenin) en modèle de Norton ce qui donne une source de courant, de courant de court-circuit $i_0 = \underline{e}/(jL\omega)$ avec une bobine d'inductance L en parallèle. On obtient alors le circuit de la figure (S16.1).

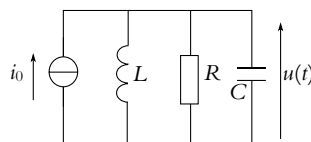


Figure S16.1

On détermine l'admittance des trois dipôles en parallèle :

$$Y = \frac{1}{jL\omega} + \frac{1}{R} + jC\omega$$

La tension $\underline{u}$ est alors $\underline{u} = i_0/Y$. On retrouve alors le résultat précédent.

2. On pose $\underline{u} = U\sqrt{2}e^{j(\omega t + \varphi)}$. Avec l'expression précédente, on obtient :

$$U = \frac{E}{\sqrt{(1 - LC\omega^2)^2 + \left(\frac{L\omega}{R}\right)^2}} = \frac{E}{\sqrt{f(\omega)}}$$

Étant donné que le numérateur est constant, U a les variations inverse de $f(\omega)$: on étudie donc cette dernière fonction.

$$\frac{df}{d\omega} = \frac{d}{d\omega} \left(1 + (LC)^2 \omega^4 - 2LC\omega^2 + \frac{L^2 \omega^2}{R^2} \right) = 4(LC)^2 \omega^3 - 4LC\omega + 2\frac{L^2}{R^2} \omega$$

La dérivée de f est nulle pour $\omega = 0$ sans condition et pour $\omega_r = \frac{1}{\sqrt{LC}} \sqrt{1 - \frac{L}{2R^2C}}$ si $R > \sqrt{\frac{L}{2C}}$.

Dans le cas $R > \sqrt{\frac{L}{2C}}$, on peut vérifier, avec le signe de la dérivée, que $\omega = 0$ correspond à un maximum pour $f(\omega)$ donc à un minimum pour U tandis que ω_r correspond à un minimum pour $f(\omega)$ donc à un maximum pour U .

Dans le cas $R < \sqrt{\frac{L}{2C}}$, on peut vérifier, avec le signe de la dérivée, que $\omega = 0$ correspond à un minimum pour $f(\omega)$ donc à un maximum pour U .

Dans les deux cas, $U(\omega = 0) = E$ et $U(\omega \rightarrow \infty) = 0$.

Comme E est un réel positif, la phase φ est telle que :

$$\varphi = \text{Arg} \left(\frac{1}{1 - LC\omega^2 + \frac{jL\omega}{R}} \right) = -\text{Arg} \left(1 - LC\omega^2 + \frac{jL\omega}{R} \right)$$

On en déduit $\tan \varphi = -\frac{L\omega}{R} \frac{1}{1 - LC\omega^2}$. Il est plus intéressant ici de considérer le signe de la partie imaginaire qui est constant et négatif, donc $\sin \varphi < 0$ et $\varphi \in [-\pi, 0]$.

On calcule la dérivée :

$$\frac{d \tan \varphi}{d\omega} = -\frac{L}{R} \frac{1 - LC\omega^2 + 2LC\omega^2}{(1 - LC\omega^2)^2} = -\frac{L}{R} \frac{1 + LC\omega^2}{(1 - LC\omega^2)^2} < 0$$

$\tan \varphi$ est donc une fonction décroissante de ω , or sur $[-\pi, -\frac{\pi}{2}[$ et $]-\frac{\pi}{2}, 0]$ la fonction $\tan$ est croissante ; on en déduit que φ est une fonction décroissante de ω sur $[-\pi, 0]$. On peut examiner les limites :

- Pour $\omega = 0$, $\tan \varphi \rightarrow 0^-$ et donc $\varphi \rightarrow 0$.
 - Pour $\omega = 1/\sqrt{LC}$, $\tan \varphi \rightarrow \pm\infty$ et $\varphi = -\pi/2$.
 - Pour $\omega \rightarrow +\infty$, $\tan \varphi \rightarrow 0^+$ et donc $\varphi \rightarrow -\pi$.
- On a tracé $U(\omega)$ sur la figure (S16.2) et φ sur la figure (S16.3) pour deux cas :

- Cas (1) $R > \sqrt{\frac{L}{2C}}$.
- Cas (2) $R < \sqrt{\frac{L}{2C}}$.

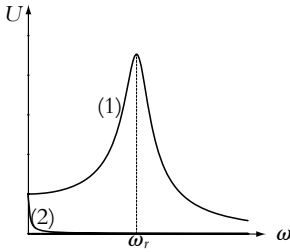


Figure S16.2

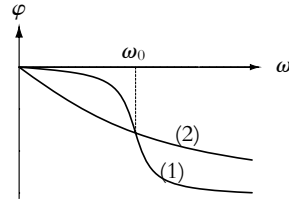


Figure S16.3

B.2 1. Par les lois d'association des impédances en série et en parallèle, on obtient :

$$\underline{Z} = \frac{R + jL\omega}{(1 - LC\omega^2) + jRC\omega}$$

2. En remarquant que $\frac{L}{R} = \frac{Q}{\omega_0}$ et $C = \frac{L}{Q^2 R^2}$, on trouve par le changement de variables proposé :

$$\underline{Z} = R \frac{1 + jQx}{1 - x^2 + j\frac{x}{Q}}$$

3. Le module de $\underline{Z}$ est : $|\underline{Z}| = R \sqrt{\frac{1 + Q^2 x^2}{x^2 + (1 - x^2)^2}}$. On pose $u = x^2$ et on cherche quand la dérivée de

$f(u) = \frac{1 + Q^2 u}{\frac{u}{Q^2} + (1 - u)^2}$ s'annule. Le calcul de la dérivée donne :

$$f'(u) = Q^2 \frac{-Q^4 u^2 - 2Q^2 u + Q^4 + 2Q^2 - 1}{(u + Q^2 (1 - u)^2)^2}$$

On cherche donc les racines de $-Q^4 u^2 - 2Q^2 u + Q^4 + 2Q^2 - 1 = 0$ dont le discriminant réduit $\Delta' = Q^4 (Q^4 + 2Q^2)$ est positif. Les racines sont :

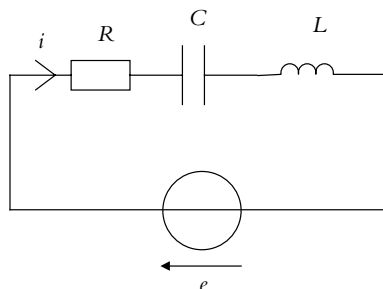
$$u = \frac{-Q^2 \pm Q^2 \sqrt{Q^4 + 2Q^2}}{Q^4}$$

Les seules solutions acceptables sont les solutions positives (on rappelle que $u = x^2$), ce qui exclut $-\frac{Q^2+Q^2\sqrt{Q^4+2Q^2}}{Q^4}$. Pour l'autre solution, il faut vérifier $1 < Q^4 + 2Q^2$ soit $Q^4 + 2Q^2 - 1 > 0$. La résolution de $g(X) = X^2 + 2X - 1 = 0$ donne $X_1 = -1 + \sqrt{2} > 0$ et $X_2 = -1 - \sqrt{2} < 0$. Comme $g(0) = -1 < 0$ et comme $X > 0$, on en déduit la condition $Q^2 > -1 + \sqrt{2}$ pour avoir un extremum.

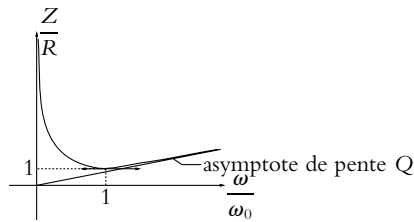
4. La pulsation à laquelle se produit l'extremum est : $\omega_r = \omega_0 \sqrt{-\frac{1}{Q^2} + \sqrt{1 + \frac{2}{Q^2}}}$.
5. Le module de l'impédance est par définition une quantité positive. D'autre part, quand ω tend vers 0 c'est-à-dire x tend vers 0, le module de l'impédance tend vers R et quand ω tend vers $+\infty$ c'est-à-dire x tend vers $+\infty$, le module de l'impédance tend vers 0. On en déduit donc que :
 - si $Q < \sqrt{-1 + \sqrt{2}}$, le module de l'impédance décroît de R à 0,
 - si $Q > \sqrt{-1 + \sqrt{2}}$, le module de l'impédance croît de R à une valeur maximale obtenue pour $\omega = \omega_r$ dont on a donné l'expression à la question précédente puis décroît jusqu'à 0.
6. D'après le schéma du circuit, on a : $\underline{\varepsilon} = \underline{Z}i$ soit en module : $E = |\underline{Z}|I$ en notant I le module de l'intensité. On en déduit que les variations de I sont opposées à celles du module de l'impédance qui vient d'être étudiée et que, par conséquent, on a un phénomène d'antirésonance pour l'intensité parcourant le générateur.
7. Si Q est grand, on peut chercher la limite de ω_r quand Q tend vers $+\infty$ soit $\omega_r \rightarrow \omega_0$. Dans ces conditions, x vaut 1 et on en déduit que :

$$|\underline{Z}| = R\sqrt{Q^2 \frac{1+Q^2}{1+Q^2(1-1)^2}} \simeq RQ^2$$

B.3 1. Le circuit considéré est le suivant :

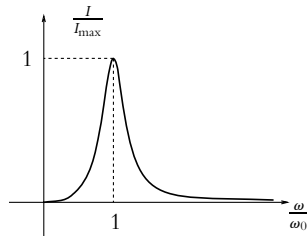


- a) L'impédance du circuit s'écrit $\underline{Z}' = R + j\left(L\omega - \frac{1}{C\omega}\right)$.
- b) Son module vaut $Z' = \sqrt{R^2 + \left(L\omega - \frac{1}{C\omega}\right)^2} = R\sqrt{1 + Q^2\left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega}\right)^2}$.
- c) La loi des mailles s'écrit $\underline{\varepsilon} = \underline{Z}'i$. Le déphasage cherché est l'opposé de la phase de $\underline{Z}'$ soit $\varphi = -\text{Arctan}Q\left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega}\right) \in \left[-\frac{\pi}{2}, \frac{\pi}{2}\right]$.
- d) On remarque que pour toute valeur de la pulsation ω , on a $\frac{Z}{R} \geq 1$. Le cas d'égalité $\frac{Z}{R} = 1$ est obtenu pour $\omega = \omega_0$ qui correspond donc à un minimum. On en déduit que $\frac{Z}{R}$ décroît de $+\infty$ à 1 pour $\omega = \omega_0$ et croît ensuite vers $+\infty$. Le comportement à l'infini est la même que celle de $Q\frac{\omega}{\omega_0}$: on a donc une asymptote oblique.

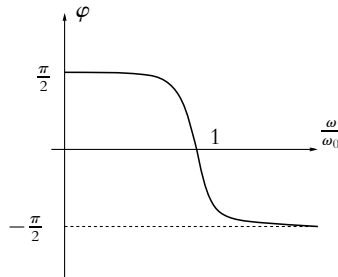


- e) L'intensité s'écrit $I = \frac{E}{Z}$ donc les variations sont inverses de celles de Z : I passe par un maximum pour $\omega = \omega_0$ et ce maximum vaut $\frac{E}{R}$.

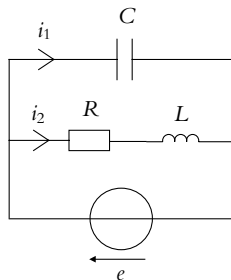
f)



- g) La fonction $f(\omega) = -L\omega + \frac{1}{C\omega}$ admet pour dérivée $f'(\omega) = -L - \frac{1}{C\omega^2} < 0$ donc la fonction f est décroissante comme $\varphi = \tan f(\omega)$. Les limites sont $\frac{\pi}{2}$ en 0 et $-\frac{\pi}{2}$ en $+\infty$. On en déduit l'allure suivante :



- h) On calcule la bande passante comme dans le cas de la résonance en intensité et on obtient $A = \frac{1}{Q}$.
- i) La quantité Q est le facteur de qualité et ω_0 la pulsation propre ou pulsation de résonance.
- j) Du fait que la fonction φ est décroissante, il suffit calculer sa valeur pour ω_1 et ω_2 pour obtenir l'intervalle demandé soit $\left[-\frac{\pi}{4}, +\frac{\pi}{4}\right]$.
2. Le nouveau circuit est le suivant :



- a) L'impédance s'écrit $\underline{Z} = \frac{R + jL\omega}{1 + jRC\omega - LC\omega^2}$.
- b) On peut mettre cette impédance sous la forme $\underline{Z} = \frac{R \left(1 + jQ \frac{\omega}{\omega_0} \right)}{jC\omega \underline{Z}'}$.
- c) Dans le cas où $Q \gg 1$, on obtient $\underline{Z} \simeq \frac{R^2 Q^2}{\underline{Z}'}$ en négligeant 1 devant $jQ \frac{\omega}{\omega_0}$ et en utilisant $L\omega_0 = \frac{1}{C\omega_0}$.
- d) $\underline{Z}'$ tend vers R pour ω tendant vers ω_0 donc $\underline{Z}$ tend vers RQ^2 .
- e) L'impédance est équivalente à une résistance comme dans le cas de l'association en série mais la valeur a changé.
- f) Le courant demandé s'écrit $\underline{i}_1 = jC\omega \underline{e}$ soit après calculs $i_1 = -\frac{E}{RQ} \sin \omega t$.
- g) Par la loi des nœuds, on en déduit $\underline{i}_2 = \underline{i} - \underline{i}_1$ et $i_2 = \frac{E}{RQ} \sqrt{1 + \frac{1}{Q^2}} \cos(\omega_0 t - \text{Arctan} Q)$.

B.4 1. a) L'impédance de l'association en série de R , L et C s'écrit $\underline{Z} = R + j \left(L\omega - \frac{1}{C\omega} \right)$.

- b) En appliquant la relation du pont diviseur de tension, on obtient

$$\underline{H} = \frac{R}{\underline{Z}} = \frac{R}{R + j \left(L\omega - \frac{1}{C\omega} \right)}$$

- c) En introduisant les quantités proposées, on obtient $\underline{H} = \frac{1}{1 + jQ \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)}$.

- d) Le gain s'obtient en prenant le module de la fonction de transfert soit $G = \frac{1}{\sqrt{1 + Q^2 \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)^2}}$.

Or $1 + Q^2 \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)^2 \geq 1$ et on a le cas d'égalité pour $\omega = \omega_0$. Dans ce cas, le gain G est maximal et sa valeur est $G_{\max} = 1$.

- e) Le déphasage est nul lorsque la partie imaginaire est nulle soit $\omega = \omega_0$.

- f) La bande passante est définie par les valeurs de pulsations ω telles que $G \geq \frac{G_{\max}}{\sqrt{2}}$ soit en appliquant la même méthode que pour la bande passante de la résonance en intensité, on obtient $-\omega_0 + \sqrt{\omega_0^2 + 4Q^2\omega_0^2} \leq \omega \leq \frac{\omega_0 + \sqrt{\omega_0^2 + 4Q^2\omega_0^2}}{2Q}$. La largeur de la bande passante s'écrit

$$\Delta\omega = \frac{\omega_0}{Q}.$$

- g) Le filtre est sélectif si la bande passante $\Delta\omega$ est petite ou encore si le facteur de qualité Q est grand et R petit à L et C fixées.
- h) Si les facteurs de qualité vérifient $Q_1 < Q_2$, les bandes passantes correspondantes sont telles que $\Delta\omega_1 > \Delta\omega_2$.
2. a) La période lue sur le graphique fourni est $T = 4$ ms, on en déduit la pulsation $\omega = \frac{2\pi}{T} = 1,57 \cdot 10^3$ rad.s⁻¹. L'amplitude de l'intensité est $I_M = 0,2$ A, celle de la tension $U_M = 8$ V et on en déduit $Z = \frac{U_M}{I_M} = 40 \Omega$.
- b) La voie I est en retard sur la voie II.
- c) Le déphasage lu sur la courbe est $\varphi = 2\pi \frac{2,5}{20} = \frac{\pi}{4}$.

d) A partir de la valeur de Z , on en déduit

$$L = \frac{1}{\omega} \left(\frac{1}{C\omega} + \sqrt{Z^2 - R^2} \right) = 62,6 \text{ mH}$$

d'après la relation établie dans la première partie. On a alors $\varphi = \text{Arctan} \frac{L\omega - \frac{1}{C\omega}}{R}$ soit numériquement $\varphi = 1,73 \neq 0,79$, ce qui est incohérent.

e) On remplace R par $R + r$ et on obtient $Z = \sqrt{(R + r)^2 + \left(L\omega - \frac{1}{C\omega} \right)^2}$ et $\tan \varphi = \frac{L\omega - \frac{1}{C\omega}}{R + r}$.

On en déduit l'expression $r = \frac{Z}{\sqrt{1 + \tan^2 \varphi}} - R$ soit numériquement $r = 8,3 \Omega$.

f) On obtient aussi $L \frac{1}{\omega} \left(\frac{1}{C\omega} + (R + r) \tan \varphi \right) = 58,6 \text{ mH}$.

B.5 1. D'après les indications de l'énoncé, on obtient le circuit de la figure (S16.4).

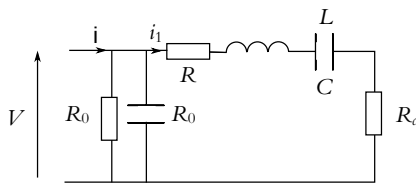


Figure S16.4

2. On s'intéresse donc à un circuit RLC série constitué d'une résistance $R + R_c$, d'une bobine L et d'un condensateur C . La loi d'Ohm donne pour les valeurs efficaces $I_1 = V/Z$ où Z est le module de $\underline{Z} = R + R_c + j \left(L\omega - \frac{1}{C\omega} \right)$. Puisque V est constante, l'intensité I est maximale si

$$Z = \sqrt{(R + R_c)^2 + \left(L\omega - \frac{1}{C\omega} \right)^2} \text{ est minimale donc si } \left(L\omega - \frac{1}{C\omega} \right) = 0 \text{ ou } \omega = 1/\sqrt{LC}.$$

3. La résonance en courant correspond à une impédance totale minimale donc une admittance totale maximale. La figure donne Y_p maximale pour x légèrement inférieur à 1 soit ω légèrement inférieure à ω_s .

L'application numérique donne $\omega = 36\,000 \text{ rad.s}^{-1}$ soit $f = \omega/2\pi = 5,6 \text{ kHz}$.

4. L'admittance $\underline{Y}$ est formée de R_0 en parallèle avec C_0 soit $Y = \sqrt{\frac{1}{R_0^2} + (C_0\omega)^2}$. Pour $\omega = \omega_s$, la valeur de l'admittance Y est $2,9 \cdot 10^{-4} \text{ S}$. D'après les questions précédentes, lorsque $\omega = \omega_s$, l'impédance $\underline{Z}_m + \underline{Z}_c$ est égale à $R + R_c$, donc $Y' = 1/(R + R_c)$ et $Y' = 0,01 \text{ S}$. On remarque donc que $Y \ll Y'$, or l'admittance totale est $\underline{Y}_{\text{tot}} = \underline{Y} + \underline{Y}'$ donc $\underline{Y}_{\text{tot}} \simeq \underline{Y}'$. La fréquence de résonance est quasiment celle de l'admittance $\underline{Y}'$.

5. Si $R_c \simeq 10 \Omega$, alors $Y' \simeq 1,6 \cdot 10^{-2} \text{ S}$. L'approximation précédente est encore mieux vérifiée, donc la fréquence de résonance est quasiment inchangée.

Chapitre 17

A.1 Lorsqu'il y a juste le chauffage, la puissance vaut $P = RI_{\text{eff}}^2$ puisque le dipôle est une résistance R . On en déduit $R = \frac{P}{I_{\text{eff}}^2} = 48 \, \Omega$.

Quand on a uniquement le moteur de la climatisation, on a $|R' + jL'\omega| = \frac{P}{I_{\text{eff}}^2} = 24 \, \Omega$.

Quand les deux sont branchés en parallèle, l'impédance s'écrit $\underline{Z} = \frac{R(R' + jL'\omega)}{R + R' + jL'\omega}$ et au niveau de la puissance $\frac{R|R' + jL'\omega|}{|R + R' + jL'\omega|} = \frac{P}{I_{\text{eff}}^2} = 20 \, \Omega$.

Des deux relations précédentes, on tire $|R + R' + jL'\omega| = \frac{48,24}{20} = 57,6 \, \Omega$ soit $(R + R')^2 + (L'\omega)^2 = R^2 + 2RR' + R'^2 + (L'\omega)^2 = 48^2 + 2RR' + 24^2 = 57,6^2$. On en déduit $R' = 4,56 \, \Omega$ puis $(L'\omega)^2 = 24^2 - 4,56^2$ à 50 Hz soit $L' = 75 \, \text{mH}$.

A.2 1. La puissance s'écrit $P = U_{\text{eff}} I_{\text{eff}} \cos \varphi$ soit un facteur de puissance

$$\cos \varphi = \frac{P}{U_{\text{eff}} I_{\text{eff}}} = 0,66$$

2. L'impédance s'écrit $\underline{Z} = R + jL\omega$ dont le module vaut

$$Z = \sqrt{R^2 + (L\omega)^2} = \frac{U_{\text{eff}}}{I_{\text{eff}}} = 19 \, \Omega$$

et le déphasage φ vérifie $\tan \varphi = \frac{L\omega}{R} = 1,14$.

On en déduit $R = \frac{Z}{\sqrt{1 + \tan^2 \varphi}} = 12,5 \, \Omega$ et $L = \frac{R \tan \varphi}{\omega} = 45,3 \, \text{mH}$.

3. Le fournisseur d'électricité souhaite limiter les pertes en ligne. Or la puissance due à ces pertes s'écrit $P_l = R_l I_{\text{eff}}^2 = R_l \frac{P^2}{U_{\text{eff}}^2 \cos^2 \varphi}$ en notant R_l la résistance de la ligne. Une des solutions pour réduire la valeur de P_l consiste à augmenter la valeur du facteur de puissance.

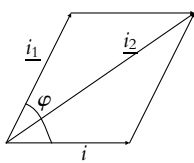
4. On place un composant en parallèle de l'installation et on souhaite que l'intensité alimentant l'ensemble et la tension à ses bornes soient en phase. Cela nécessite que l'impédance de l'association en parallèle soit réelle ou encore son argument nul. Comme l'inductance introduit un déphasage positif, il faut introduire un déphasage négatif, ce qui est possible avec une capacité.

5. L'impédance de l'association en parallèle s'écrit dans ces conditions

$$\frac{R + jL\omega}{(1 - LC\omega^2) + jRC\omega} = \frac{R + j\omega(L - L^2 C \omega^2 - R^2 C)}{(1 - LC\omega^2)^2 + (RC\omega)^2}$$

La condition pour avoir un facteur de puissance égal à 1, il faut $L = C(L^2 \omega^2 + R^2)$ donc $C = 126 \, \mu\text{F}$.

A.3 On peut représenter les intensités vectoriellement soit



On a $\underline{i}_2 = \underline{i} + \underline{i}_1$ soit en norme $I_2^2 = I^2 + I_1^2 + 2II_1 \cos(2\pi - \varphi)$ dont on déduit $I \cos \varphi = \frac{I^2 + I_1^2 - I_2^2}{2I_1}$.

La puissance pour l'appareil s'écrit

$$P = UI \cos \varphi = \frac{R}{2} (I^2 + I_1^2 - I_2^2)$$

A.4 On note φ' l'argument de l'impédance totale $R + \underline{Z}$ et φ celui de $\underline{Z}$. On prend l'origine des phases pour l'intensité $\underline{i}$, soit :

- $\underline{u}_1 = U_1 \sqrt{2} \exp(j\omega t)$,
- $\underline{u}_2 = U_2 \sqrt{2} \exp(j(\omega t + \varphi))$,
- $\underline{u} = U \sqrt{2} \exp(j(\omega t + \varphi'))$.

La puissance reçue par $\underline{Z}$ est $\mathcal{P} = U_2 I \cos \varphi = \frac{U_2 U_1}{R} \cos \varphi$.

Il reste donc à déterminer $\cos \varphi$.

La loi d'additivité des tensions donne :

$$\underline{u} = \underline{u}_1 + \underline{u}_2$$

soit :

$$U \sqrt{2} e^{j(\omega t + \varphi')} = U_1 \sqrt{2} e^{j\omega t} + U_2 \sqrt{2} e^{j(\omega t + \varphi)}$$

Si on multiplie par la quantité complexe conjuguée, on trouve :

$$U^2 = U_1^2 + U_2^2 + 2U_1 U_2 \cos \varphi \quad \text{d'où} : \quad \mathcal{P} = \frac{U_2 U_1}{R} \left(\frac{U^2 - U_1^2 - U_2^2}{2U_1 U_2} \right) = \frac{U^2 - U_1^2 - U_2^2}{2R}$$

On peut également procéder par construction de Fresnel comme dans l'exercice précédent.

A.5 1. On calcule tout d'abord l'impédance du dipôle équivalent au circuit soit :

- Pour R et C en parallèle : $\underline{Y} = jC\omega + \frac{1}{R}$ soit une impédance $\underline{Z} = \frac{R}{1 + jRC\omega}$

- Cette impédance est en série avec L ce qui donne : $\underline{Z}_{\text{eq}} = jL\omega + \underline{Z}$

Pour que l'impédance totale soit équivalente à une résistance, il faut qu'elle soit réelle. On exprime donc la partie imaginaire de $\underline{Z}_{\text{eq}}$:

$$\underline{Z}_{\text{eq}} = jL\omega + \frac{R}{1 + (RC\omega)^2} (1 - jRC\omega) \Rightarrow \text{Im}(\underline{Z}_{\text{eq}}) = L\omega - \frac{R^2 C \omega}{1 + (RC\omega)^2}$$

Pour avoir $\underline{Z}_{\text{eq}}$ réelle il faut choisir : $L = \frac{R^2 C}{1 + (RC\omega)^2}$ et donc $\underline{Z}_{\text{eq}} = R_{\text{eq}} = \frac{R}{1 + (RC\omega)^2}$

2. L'application numérique donne $L = 0,12 \text{ H}$.

3. Si l'on note $E = E_0 / \sqrt{2}$ la valeur efficace de la f.e.m. du générateur, on peut écrire $I = E / R_{\text{eq}}$ soit $I = \frac{E}{R} (1 + R^2 C^2 \omega^2)$. L'application numérique donne $I = 5 \text{ A}$.

4. La tension efficace aux bornes de L est $U_{AD} = L\omega I = 240 \text{ V}$. Pour calculer U_{DB} , il faut calculer le module de $\underline{Z}$ soit

$$U_{DB} = \left| \frac{R}{1 + jRC\omega} \right| I = \frac{R}{\sqrt{1 + (RC\omega)^2}} I = 300 \text{ V}$$

On constate qu'il ne faut pas ajouter les valeurs efficaces : $E \neq U_{AD} + U_{DB}$.

5. On pose $i_1 = I_1 \sqrt{2} \cos(\omega t + \varphi_1)$ et $i_2 = I_2 \sqrt{2} \cos(\omega t + \varphi_2)$, les intensités réelles. Les grandeurs complexes seront soulignées comme d'habitude. L'origine des phases est celle du générateur donc de la tension u_{AB} .

On considère toujours que l'impédance totale est égale à R_{eq} . Un diviseur de tension entre u_{AB} et u_{DB} donne :

$$\underline{u}_{DB} = \underline{u}_{DB} \frac{\underline{Z}}{R_{\text{eq}}} = \underline{\varepsilon} \frac{\underline{Z}}{R_{\text{eq}}}$$

On peut ainsi écrire les relations pour les intensités, d'abord pour $i_1 = \underline{u}_{DB}/R$:

$$i_1 = \underline{e} \frac{R}{1 + jRC\omega} \frac{1 + (RC\omega)^2}{R} \frac{1}{R} = \frac{E_0}{R} (1 - jRC\omega)$$

On en déduit la valeur efficace $I_1 = E\sqrt{1 + (RC\omega)^2}/R$, soit $I_1 = 3$ A. En ce qui concerne la phase, on peut écrire $\tan \varphi_1 = -RC\omega$, or la partie réelle de i_1 est positive, donc $\cos \varphi_1$ aussi et $\varphi_1 \in \left[-\frac{\pi}{2}, \frac{\pi}{2}\right]$. Finalement on a $\varphi_1 = -\arctan(RC\omega) = -0,927$ rad. En ce qui concerne $i_2 = jC\omega \underline{u}_{DB}$:

$$i_2 = jC\omega \underline{e} \frac{R}{1 + jRC\omega} \frac{1 + (RC\omega)^2}{R} = E_0 jC\omega (1 - jRC\omega) = jRC\omega i_1$$

On en déduit $I_2 = RC\omega I_1$ soit $I_2 = 4$ A et $\varphi_2 = \varphi_1 + \frac{\pi}{2}$.

6. La puissance moyenne $\langle P \rangle$ est égale à $R_{eq} i^2$ soit 900 W.

B.1 1. On adopte les notations de la figure (S17.1).

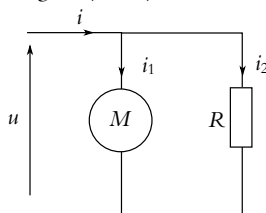


Figure S17.1

La valeur moyenne de la puissance reçue par l'ensemble de l'installation est la somme de celle reçue par le moteur et celle reçue par les lampes soit $P = P_1 + P_2$ or $P_2 = U^2/R = 576$ W, d'où $P = 4\,576$ W.

La puissance reçue par le moteur est $P_1 = UI_1 \cos \varphi_1$, ce qui permet de déterminer l'intensité efficace traversant le moteur : $I_1 = 20,8$ A.

En ce qui concerne, l'intensité efficace traversant les lampes, il suffit d'écrire $I_2 = U/R = 2,4$ A.

2. a) Attention, il ne faut surtout pas additionner les valeurs efficaces des intensités, mais les valeurs instantanées. Ainsi, on peut écrire : $i = i_1 + i_2$ ou

$$I\sqrt{2}e^{j(\omega t - \varphi)} = I_1\sqrt{2}e^{j(\omega t - \varphi_1)} + I_2\sqrt{2}e^{j(\omega t - \varphi_2)}$$

où φ , φ_1 et φ_2 sont les déphasages des intensités par rapport à $\underline{u}$ pour chaque branche. La branche dans laquelle circule i_2 est constituée de résistors donc $\varphi_2 = 0$. On obtient, en simplifiant l'expression :

$$Ie^{-j\varphi} = I_1e^{-j\varphi_1} + I_2$$

On multiplie cette expression par son complexe conjugué, ce qui revient à calculer le module au carré : $I^2 = I_1^2 + I_2^2 + 2I_1I_2 \cos \varphi_1$. On exprime ensuite les puissances moyennes : $P = UI \cos \varphi$, $P_1 = UI_1 \cos \varphi_1$ et $P_2 = UI_2$ ce qui entraîne :

$$\frac{P^2}{U^2 \cos^2 \varphi} = \frac{P_1^2}{U^2 \cos^2 \varphi_1} + \frac{P_2^2}{U^2} + 2 \frac{P_1 P_2}{U^2}$$

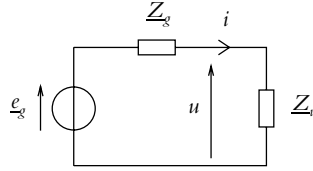
b) Dans l'expression précédente, on peut simplifier par U .

Sachant que $1 + \tan^2 \varphi_1 = 1/\cos^2 \varphi_1$, on a :

$$\frac{P^2}{\cos^2 \varphi} = P_1^2 (1 + \tan^2 \varphi_1) + P_2^2 + 2P_1 P_2 = (P_1 + P_2)^2 + P_1^2 \tan^2 \varphi_1$$

En écrivant que $P = P_1 + P_2$, en inversant l'expression et en prenant la racine carrée, on arrive à l'expression demandée.

B.2 1. Le circuit étudié est le suivant :



On note $\underline{Z}_g = R_g + jX_g$, $\underline{Z}_u = R_u + jX_u$, $e_g = E\sqrt{2} \exp(j\omega t)$ et $i = I\sqrt{2} \exp(j(\omega t - \varphi))$. D'après le cours, la puissance reçue par l'impédance $\underline{Z}_u$ est : $\mathcal{P}_m = R_u I^2$. Or une loi des mailles donne : $i = \frac{e_g}{(R_g + R_u) + j(X_g + X_u)}$ soit en prenant le module : $I = \frac{E}{\sqrt{(R_g + R_u)^2 + (X_g + X_u)^2}}$. La puissance s'écrit :

$$\mathcal{P}_m = \frac{R_u E^2}{(R_g + R_u)^2 + (X_g + X_u)^2} \quad (1)$$

On cherche donc les valeurs de R_u et X_u qui rendent la puissance maximale. Pour cela, il faut tout d'abord calculer les dérivées de $\mathcal{P}_m$ par rapport à R_u et X_u et déterminer à quelles conditions elles sont nulles.

Ces dérivées s'écrivent :

$$\begin{cases} \frac{\partial \mathcal{P}_m}{\partial X_u} = \frac{-2R_u E^2 (X_g + X_u)}{\left((R_g + R_u)^2 + (X_g + X_u)^2\right)^2} \\ \frac{\partial \mathcal{P}_m}{\partial R_u} = E^2 \frac{(R_g + R_u)(R_g - R_u) + (X_g + X_u)^2}{\left((R_g + R_u)^2 + (X_g + X_u)^2\right)^2} \end{cases}$$

L'annulation de $\frac{\partial \mathcal{P}_m}{\partial X_u}$ conduit à prendre $X_u = -X_g$ et en reportant dans l'annulation de $\frac{\partial \mathcal{P}_m}{\partial R_u}$, on obtient : $R_u = R_g$.

Pour être sûr d'avoir un maximum, on étudie les valeurs limites de $\mathcal{P}_m$ quand R_u et X_u tendent vers 0 ou $+\infty$ comme cela a été fait à l'exercice précédent. On a alors un maximum de puissance transférée à la charge d'impédance $\underline{Z}_u$: on dit qu'on a réalisé une adaptation d'impédance au transfert de puissance ou que la charge est adaptée au transfert de puissance.

2. Pour avoir un transfert maximal de puissance, il faut d'après la question précédente $\underline{Z}_g = \underline{Z}_u^*$ soit $R_g = R_u$ et $X_g = -X_u$. Or ici l'impédance interne du générateur est réelle donc il faut que la partie imaginaire de l'admittance du dipôle d'utilisation soit nulle. On calcule l'impédance de ce dipôle :

$$\underline{Z}_u = \frac{1}{jC\omega} + \frac{jL\omega R}{R + jL\omega} = \frac{1}{jC\omega} + \frac{jL\omega R (R - jL\omega)}{R^2 + L^2\omega^2}$$

Pour obtenir l'adaptation, on doit avoir :

$$\begin{cases} \frac{1}{C\omega} = \frac{L\omega R^2}{R^2 + L^2\omega^2} \\ \frac{R(L\omega)^2}{R^2 + (L\omega)^2} = R_g \end{cases} \quad (2)$$

On tire de ces deux équations :

$$\begin{cases} L\omega = \sqrt{\frac{R_g R^2}{R - R_g}} \\ C\omega = \sqrt{\frac{1}{(R - R_g) R_g}} \end{cases} \quad (3)$$

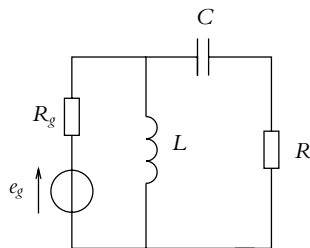
3. D'après les expressions précédentes, on voit qu'il faut $R > R_g$.

Si on place L avant C , l'admittance $\underline{Y}_u$ est alors :

$$\underline{Y}_u = \frac{1}{jL\omega} + \frac{jC\omega}{1 + jRC\omega}$$

soit :

$$\underline{Y}_u = \frac{1}{jL\omega} + \frac{jC\omega(1 - jRC\omega)}{1 + (RC\omega)^2}$$



Pour obtenir l'adaptation d'impédance, il faut que la partie réelle de $\underline{Y}_u$ soit égale à $\frac{1}{R_g}$ et que sa partie imaginaire soit nulle :

$$\begin{cases} \frac{1}{L\omega} = \frac{C\omega}{1 + (RC\omega)^2} \\ \frac{(C\omega)^2 R}{1 + (RC\omega)^2} = \frac{1}{R_g} \end{cases} \quad (4)$$

On tire de ces deux équations :

$$\begin{cases} L\omega = \sqrt{\frac{R_g^2 R}{R_g - R}} \\ C\omega = \sqrt{\frac{1}{(R_g - R)R}} \end{cases} \quad (5)$$

Cette solution est valable lorsque $R_g > R$.

REMARQUE : dans cet exercice, on a utilisé la première fois l'impédance et la deuxième fois l'admittance. Le lecteur pourra s'entraîner en faisant l'autre choix.

- B.3** 1. a) On sait que la puissance moyenne est $P_m = IU \cos \varphi$ donc l'intensité efficace est $I = P_m / (U \cos \varphi) = 30,5 \text{ A}$.

- b) On peut écrire d'autre part $P_m = RI^2$ car seule la partie résistive consomme de l'énergie. Ainsi $R = \frac{U^2 \cos^2 \varphi}{P_m}$ soit $R = 4,7 \Omega$.

L'impédance du moteur est $\underline{Z} = R + jL\omega$, donc on peut écrire $\tan \varphi = L\omega/R$ soit $L = \frac{R \tan \varphi}{2\pi f}$.

L'application numérique donne $L = 20 \text{ mH}$.

2. a) On calcule l'admittance équivalente à l'ensemble du condensateur en parallèle sur (L, R) :

$$\underline{Y}_{\text{eq}} = jC\omega + \frac{1}{R + jL\omega} = \frac{R}{R^2 + (L\omega)^2} + j \left(C\omega - \frac{L\omega}{R^2 + (L\omega)^2} \right)$$

On note φ' , l'argument de $\underline{Z}_{\text{eq}} = 1/\underline{Y}_{\text{eq}}$. L'expression précédente permet de calculer :

$$\tan(-\varphi') = \frac{C\omega - \frac{L\omega}{R^2 + (L\omega)^2}}{\frac{R}{R^2 + (L\omega)^2}} = \frac{C\omega(R^2 + (L\omega)^2)}{R} - \tan \varphi$$

On obtient donc :

$$\begin{aligned} C &= \frac{R}{\omega(R^2 + (L\omega)^2)} (\tan \varphi - \tan \varphi') = \frac{RI^2}{\omega U^2} (\tan \varphi - \tan \varphi') \\ &= \frac{P_m}{2\pi f U^2} (\tan \varphi - \tan \varphi') \end{aligned}$$

À une valeur de $\cos \varphi'$ correspondent deux valeurs de $\tan \varphi'$, l'une positive et l'autre négative, ce qui donne deux valeurs de C possibles. Pour avoir C la plus petite possible, on choisit $\tan \varphi' > 0$ soit $\tan \varphi' = 0,48$. On trouve $C = 207 \mu\text{F}$.

- b) La puissance absorbée par l'installation est toujours la même puisqu'en moyenne il n'y a pas de puissance absorbée par C .
- c) L'expression de l'intensité est $I' = P_m / (U \cos \varphi')$ soit $I' = 20,4$ A au lieu de 30,5 A. La consommation d'énergie par effet Joule dans les câbles de l'installation est donc moins importante. c'est le but recherché.

B.4 1. En notation complexe on peut écrire : $\underline{u} = U\sqrt{2}e^{j\omega t}$, $\underline{i} = I\sqrt{2}e^{j(\omega t + \varphi)}$. La loi d'Ohm donne $\underline{u} = \underline{Z}\underline{i}$ d'où $\underline{Z} = \frac{U}{I}e^{-j\varphi}$ et $Z = \frac{U}{I}$.

- 2. a) Pour calculer la puissance, il faut utiliser les grandeurs réelles. La puissance moyenne P_e est obtenue par l'intégrale :

$$P_e = \frac{1}{T} \int_0^T u(t)i(t) dt = \frac{1}{T} \int_0^T 2UI \cos(2\pi ft) \cos(2\pi ft + \varphi) dt$$

ce qui entraîne :

$$P_e = \frac{UI}{T} \int_0^T [\cos(4\pi ft + \varphi) + \cos \varphi] dt = UI \cos \varphi$$

- b) L'expression complexe de $\underline{Z}$ permet d'écrire $\underline{Z} = \frac{U}{I}(\cos \varphi - j \sin \varphi)$. Ainsi la partie réelle de $\underline{Z}$ est $R = \frac{U}{I} \cos \varphi$. On peut donc remplacer $\cos \varphi$ dans l'expression de P_e par RI/U soit R/Z . On remplace aussi I par U/Z ce qui conduit à $P_e = RU^2/Z^2$.
- 3. a) L'impédance du dipôle D est $\underline{Z} = R + jL\omega$ d'où d'après ce qui précède :

$$P_e = \frac{RU^2}{R^2 + L^2\omega^2} \Rightarrow R^2 - \frac{U^2}{P_e}R + L^2\omega^2 = 0 \Rightarrow R = \frac{U^2}{2P_e} \pm \frac{1}{2}\sqrt{\frac{U^4}{P_e^2} - 4L^2\omega^2}$$

On trouve alors $R_1 = 18,1 \Omega$ et $R_2 = 4,9 \Omega$.

- b) On sait que seule la partie réelle (résistive) consomme de l'énergie dans un dipôle donc $P_e = RI^2$ ce qui donne les deux valeurs d'intensité efficace possibles : $I_1 = 11,7$ A et $I_2 = 22,6$ A.
- c) L'intensité admissible sur la ligne est au maximum 16 A donc seule la résistance R_1 permet de ne pas dépasser cette valeur.
- d) On a donc $I = I_1$ dans la ligne. La puissance consommée dans la ligne est donc $P_0 = R_0I_1^2 = 164$ W.
- 4. a) Si l'on suppose la tension U inchangée alors l'intensité efficace dans D est toujours I_1 . En ce qui concerne le condensateur, l'intensité efficace est $I_c = Z_cU$ ou $I_c = C\omega U$ soit $I_c = 9,832$ A.
- b) Les deux dipôles C et D sont en parallèle ; l'admittance du dipôle équivalent est donc :

$$\underline{Y} = jC\omega + \frac{1}{R_1 + jL\omega} \text{ avec } \underline{I}'_0 = \underline{Y}U$$

Si l'on passe aux valeurs efficaces :

$$I'_0 = |\underline{Y}|U = \sqrt{\left(\frac{R_1}{R_1^2 + L^2\omega^2}\right)^2 + \left(C\omega - \frac{L\omega}{R_1^2 + L^2\omega^2}\right)^2} U$$

Ce qui donne $I'_0 = 11,3$ A.

- c) La puissance P'_0 vaut $R_0I'^2_0$ soit 154 W.
- d) L'intensité dans la ligne est légèrement inférieure à celle en l'absence de condensateur. L'avantage est que pour une même puissance consommée par le dipôle D , la puissance « perdue » dans la ligne est moins importante.

Chapitre 18

A.1 Fonction de transfert d'un pont

1. La tension $\underline{V}_s$ est égale à $\underline{V}_A - \underline{V}_B$. Si l'on note $\underline{v}_1$ la tension $\underline{V}_A - \underline{V}_E$ et $\underline{v}_2$ la tension $\underline{V}_E - \underline{V}_B$, on peut écrire : $\underline{V}_s = \underline{v}_1 + \underline{v}_2$.

Comme il n'y a pas de dipôle entre A et B , les deux résistances R sont en série dans la branche EBD , et r et C sont en série dans la branche EAD . On peut utiliser la relation du diviseur de tension entre $\underline{v}_1$ et $\underline{V}_e$, et entre $\underline{v}_2$ et $\underline{V}_e$, soit :

$$\underline{v}_1 = -\frac{r}{r + \frac{1}{jC\omega}} \underline{V}_e = -\frac{j r C \omega}{1 + j r C \omega} \underline{V}_e \quad \text{et} \quad \underline{v}_2 = \frac{R}{2R} \underline{V}_e = \frac{1}{2} \underline{V}_e$$

On en déduit :

$$\underline{V}_s = \underline{v}_1 + \underline{v}_2 = \left(-\frac{j r C \omega}{1 + j r C \omega} + \frac{1}{2} \right) \underline{V}_e \Rightarrow \underline{T}(j\omega) = \frac{\underline{V}_s}{\underline{V}_e} = \frac{1}{2} \frac{1 - j r C \omega}{1 + j r C \omega}$$

Si l'on pose $T_0 = 1/2$ et $\omega_1 = 1/rC$, on peut écrire $\underline{T}$ sous la forme :

$$\underline{T} = T_0 \frac{1 - j \frac{\omega}{\omega_1}}{1 + j \frac{\omega}{\omega_1}}$$

2. Le module de $\underline{T}$ est égal à :

$$T(\omega) = |\underline{T}| = T_0 \frac{\sqrt{1 + \left(\frac{\omega}{\omega_1}\right)^2}}{\sqrt{1 + \left(\frac{\omega}{\omega_1}\right)^2}} = T_0$$

Le déphasage φ_1 est tel que :

$$\varphi_1 = \text{Arg}(\underline{T}) = \text{Arg}(T_0) + \text{Arg}(1 - j r C \omega) - \text{Arg}(1 + j r C \omega) = 0 + \varphi_n - \varphi_d$$

On peut écrire $\tan \varphi_n = -r C \omega$ et comme $\cos \varphi_n > 0$, alors $\varphi_n \in \left[-\frac{\pi}{2}, \frac{\pi}{2}\right]$. On a donc $\varphi_n = -\arctan(r C \omega)$.

De même, $\tan \varphi_d = r C \omega$ et $\cos \varphi_d > 0$ donc $\varphi_d \in \left[-\frac{\pi}{2}, \frac{\pi}{2}\right]$ et $\varphi_d = \arctan(r C \omega)$.

On en déduit que $\varphi_1 = -2 \arctan(r C \omega)$. C'est une fonction décroissante de ω . Lorsque $\omega = 0$ alors $\varphi_1 = 0$ et lorsque $\omega \rightarrow \infty$ alors $\varphi_1 \rightarrow -\pi$.

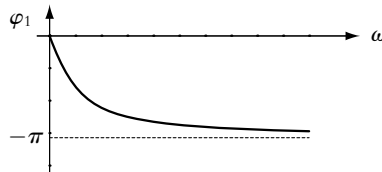


Figure S18.1

3. Pour avoir $\varphi = -\pi/2$, il faut $r_0 C \omega = 1/\sqrt{2}$ soit $r_0 = 707 \, \Omega$.

A.2 1. La réponse est 3. Comme on est en régime sinusoïdal, l'amplificateur opérationnel est en régime linéaire. Par ailleurs, on considère l'amplificateur opérationnel est idéal donc on a $\epsilon = v_+ - v_- = 0$. Or par la relation du pont diviseur de tension, on a $\underline{v}_+ = \frac{\underline{v}_B}{1 + jRC\omega}$. Par ailleurs, $\underline{v}_- = \underline{s} = \underline{v}_A$. On applique le théorème de Millman soit $\underline{v}_B = (1 + jRC\omega) \underline{s}$. On en déduit

$$\underline{H} = \frac{\underline{s}}{\underline{e}} = \frac{1}{1 + 2jRC\omega - R^2 C C' \omega^2}.$$

2. La réponse est 2. En effet, le calcul du gain donne

$$G = \frac{1}{\sqrt{1 + R^4 C^2 C'^2 \omega^4 + (4R^2 C^2 - 2R^2 C C') \omega^2}}$$

donc pour obtenir la forme souhaitée, il faut $C' = 2C$.

3. La réponse est 2 : on obtient par identification dans l'expression précédente $\omega_0 = \frac{1}{R\sqrt{CC'}} = \frac{1}{RC\sqrt{2}}$.
4. La réponse est 1 car quand ω tend vers 0, on a un gain de 1 soit $G_{dB} = 0$.
5. La réponse est 3 car pour ω tendant vers l'infini, on a l'équivalence $G \simeq \frac{1}{R^2 C C' \omega^2}$ soit $G_{dB} = 40 \log_{10} \frac{\omega_0}{\omega}$.
6. La réponse est 1 car pour ω tendant vers 0, la fonction de transfert tend vers 1 et le déphasage est nul.
7. La réponse est 2. En effet, pour ω infini, la fonction de transfert est équivalente à $\frac{1}{-R^2 C C' \omega^2}$ soit un déphasage de π .
8. Les réponses sont 1 et 2. En effet, on a un filtre actif puisque l'amplificateur opérationnel doit être alimenté et il faut lui fournir de l'énergie. D'autre part, c'est un filtre passe-bas car il ne laisse passer que les basses fréquences.

A.3 1. L'hypothèse que le filtre fonctionne à vide impose l'absence de courant en sortie du filtre où il est donc possible d'appliquer le théorème de Millman. Par ailleurs, on utilisera aussi ce théorème en A entre les deux capacités C soit $\frac{v_A}{1} = \frac{jRC\omega(\underline{e} + \underline{s})}{1 + 2jRC\omega}$ et $\underline{s} = \frac{\underline{e} + jC\omega(r + jL\omega)v_A}{1 + jrC\omega - LC\omega^2}$. On élimine $\frac{v_A}{1}$ entre les deux expressions précédentes pour obtenir la fonction de transfert $\underline{H} = \frac{(1 - rRC^2\omega^2) + jRC\omega(2 - LC\omega^2)}{1 - (L + RrC)C\omega^2 + jC\omega(r + 2R - RLC\omega^2)}$.

2. Pour obtenir un réjecteur de fréquence pour la pulsation ω_0 , il faut que la fonction de transfert $\underline{H}$ s'annule pour $\omega = \omega_0$ et $R = R_0$. Cela impose aux parties réelle et imaginaire du numérateur d'être nulles dans ces conditions. On en déduit $\omega_0 = \sqrt{\frac{2}{LC}}$ et $R_0 = \frac{L}{2rC}$.

3. On effectue le changement de variables proposé et par identification, on a $A = \frac{2x}{Q(1 - x^2)}$. On peut vérifier que les parties imaginaires des numérateur et dénominateur sont égaux.

4. On étudie la fonction $f(x) = 1 + A^2$ en fonction de x . La dérivée vaut $f'(x) = \frac{8x(1 + x^2)}{Q^2(1 - x^2)^2(1 - x^2)}$: elle est positive pour $x < 1$ et négative sinon. On en déduit que le gain est décroissant pour $x < 1$ et croissant sinon. On note que le gain prend la valeur de 1 lorsque ω est nul ou infini et une valeur de 0 pour ω_0 .

Quant au déphasage, il s'exprime par $\varphi = \text{Arctan} \frac{2x}{Q(x^2 - 1)}$. Or la dérivée de $\tan \varphi$ vaut $\frac{-2(1 + x^2)}{Q(x^2 - 1)^2}$, elle est négative donc le déphasage est décroissant de 0 à $-\frac{\pi}{2}$ puis de $\frac{\pi}{2}$ à 0.

5. La bande passante s'obtient en résolvant $\frac{1}{\sqrt{1 + A^2}} \geq \frac{1}{\sqrt{2}}$. On obtient une bande coupée entre $\frac{-r + \sqrt{r^2 + 2\frac{L}{C}}}{L}$ et $\frac{r + \sqrt{r^2 + 2\frac{L}{C}}}{L}$.

B.1 1. On nomme $\underline{Z}_1$ l'impédance formée de l'association de R_1 et de C_1 en parallèle et $\underline{Z}_2$ l'impédance formée de l'association de R_2 et de C_2 en série. La relation du pont diviseur de tension

donne :

$$\underline{H} = \frac{v_s}{v_e} = \frac{Z_2}{Z_2 + Z_1}$$

Étant donné que R_1 et C_1 sont en parallèle, il est préférable d'utiliser l'admittance $\underline{Y}_1$ au lieu de l'impédance $\underline{Z}_1$. On obtient donc :

$$\underline{H} = \frac{Z_2 \underline{Y}_1}{Z_2 \underline{Y}_1 + 1}$$

d'où :

$$\underline{H} = \frac{\left(R_2 + \frac{1}{jC_2\omega}\right) \left(jC_1\omega + \frac{1}{R_1}\right)}{1 + \left(R_2 + \frac{1}{jC_2\omega}\right) \left(jC_1\omega + \frac{1}{R_1}\right)}$$

Soit en multipliant numérateur et dénominateur par R_1 et $jC_2\omega$:

$$\underline{H}(j\omega) = \frac{f(j\omega)}{f(j\omega) + jR_1C_2\omega}$$

avec :

$$f(j\omega) = (1 + jR_2C_2\omega)(1 + jR_1C_1\omega)$$

2. Le numérateur est sous la forme factorisée demandée avec $\tau_1 = R_1C_1$ et $\tau_2 = R_2C_2$.

Pour calculer τ_3 et τ_4 , on développe le dénominateur :

$$1 + j(R_1C_2 + \tau_1 + \tau_2)\omega - \tau_1\tau_2\omega^2$$

expression qu'il faut identifier à :

$$1 + j(\tau_3 + \tau_4)\omega - \tau_3\tau_4\omega^2$$

On trouve alors :

$$\begin{cases} \tau_3 + \tau_4 = R_1C_2 + \tau_1 + \tau_2 \\ \tau_3\tau_4 = \tau_1\tau_2 \end{cases}$$

Puisque les deux expressions précédentes donnent la somme et le produit de τ_3 et τ_4 , ces deux grandeurs sont solutions de :

$$\tau^2 - (R_1C_2 + \tau_1 + \tau_2)\tau + \tau_1\tau_2 = 0$$

Dans le cas où il existe deux solutions réelles (donc quand le discriminant de l'équation du second degré en ω vérifie $\Delta > 0$, ce qui correspond à l'hypothèse de l'énoncé), elles sont positives (somme et produit positifs) et valent :

$$\frac{\tau_1 + \tau_2 + R_1C_2}{2} \pm \frac{\sqrt{(\tau_1 + \tau_2 + R_1C_2)^2 - 4\tau_1\tau_2}}{2}$$

3. Il faut tout d'abord déterminer R_1C_2 sachant que : $\frac{\tau_1}{\tau_2} = 0,1 = \frac{R_1C_1}{R_2C_2} = \frac{R_1}{R_2}\lambda$ donc $R_1C_2 = \tau_2 \frac{R_1}{R_2} = \tau_2 \frac{1}{\lambda} \frac{\tau_1}{\tau_2} = \frac{\tau_1}{\lambda}$. On en déduit avec les expressions précédentes :

$$\tau_3 = \frac{\tau}{2} \left(11 + \lambda - \sqrt{(11 + \lambda)^2 - 40} \right)$$

Si on pose $X = 11 + \lambda$ et comme $\tau_3 = \frac{\tau}{5}$, l'équation précédente revient à résoudre : $2 = 5(X - \sqrt{X^2 - 40})$ qui a pour solution $X = 50,2$ soit $\lambda = 39,2$.

Alors :

$$\tau_4 = \tau \left(\frac{11 + \lambda}{2} + \frac{\sqrt{(11 + \lambda)^2 - 40}}{2} \right) = 5,10^{-2} \text{ s}$$

4. Construction du diagramme de Bode : Finalement la fonction de transfert s'écrit sous la forme :

$$\underline{H} = \frac{(1 + j\tau_1\omega)(1 + j\tau_2\omega)}{(1 + j\tau_3\omega)(1 + j\tau_4\omega)}$$

avec les pulsations correspondantes

- $\omega_1 = \frac{1}{\tau_1} = 1\,000 \text{ rad.s}^{-1}$,
- $\omega_2 = \frac{1}{\tau_2} = 100 \text{ rad.s}^{-1}$,
- $\omega_3 = \frac{1}{\tau_3} = 5\,000 \text{ rad.s}^{-1}$,
- $\omega_4 = \frac{1}{\tau_4} = 20 \text{ rad.s}^{-1}$.

Le gain en décibel s'écrit :

$$G = G_1 + G_2 + G_3 + G_4$$

avec :

- $G_1 = 20 \log \sqrt{1 + (\tau_1 \omega)^2}$;
- $G_2 = 20 \log \sqrt{1 + (\tau_2 \omega)^2}$;
- $G_3 = -20 \log \sqrt{1 + (\tau_3 \omega)^2}$;
- $G_4 = -20 \log \sqrt{1 + (\tau_4 \omega)^2}$;

Les comportements asymptotiques sont les suivants :

- pour G_1 : $\omega \ll \omega_1$ $G_1 \simeq 0$ et $\omega \gg \omega_1$ $G_1 \simeq 20 \log \tau_1 \omega$ soit une pente de 20 dB/déc ;
- pour G_2 : $\omega \ll \omega_2$ $G_2 \simeq 0$ et $\omega \gg \omega_2$ $G_2 \simeq 20 \log \tau_2 \omega$ soit une pente de 20 dB/déc ;
- pour G_3 : $\omega \ll \omega_3$ $G_3 \simeq 0$ et $\omega \gg \omega_3$ $G_3 \simeq -20 \log \tau_3 \omega$ soit une pente de -20 dB/déc ;
- pour G_4 : $\omega \ll \omega_4$ $G_4 \simeq 0$ et $\omega \gg \omega_4$ $G_4 \simeq -20 \log \tau_4 \omega$ soit une pente de -20 dB/déc.

Les différents diagrammes ainsi que le diagramme résultant sont tracés sur la figure ci-contre avec les asymptotes :

- pour $\omega < \omega_4$: 0 dB/déc ;
- pour $\omega_4 < \omega < \omega_2$: -20 dB/déc ;
- pour $\omega_2 < \omega < \omega_1$: 0 dB/déc ;
- pour $\omega_1 < \omega < \omega_3$: +20 dB/déc ;
- pour $\omega_3 < \omega$, 0 dB/déc.

Les asymptotes des différents filtres du premier ordre sont représentées en traits pointillés. Le diagramme asymptotique du filtre est en trait gras et le diagramme réel en trait fin.

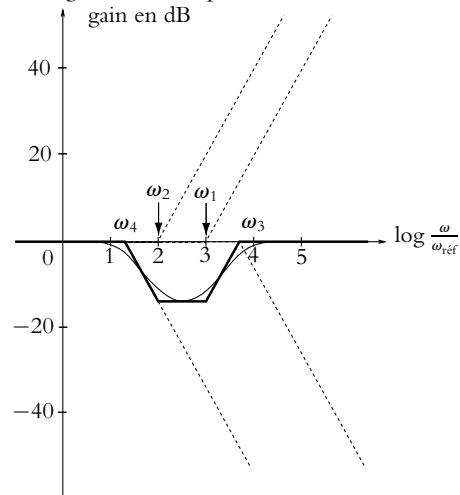


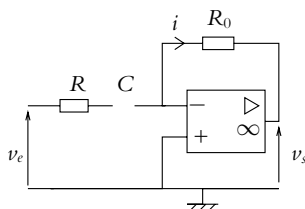
Diagramme de Bode en gain

B.2 1. Il existe un bouclage entre l'entrée inverseuse et la sortie : on peut supposer que l'amplificateur opérationnel fonctionne en régime linéaire. Comme de plus, l'amplificateur opérationnel est idéal, les courants i_+ et i_- sont nuls et on peut écrire $v_+ = v_-$.

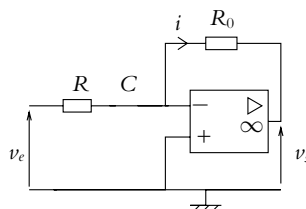
a) Détermination du type de filtre :

Dans le cas des très basses fréquences, le circuit est équivalent au circuit représenté ci-dessous. Comme $i = 0$, $v_s = v_- = v_+ = 0$ et $H = 0$.

Dans le cas des hautes fréquences, puisque $v_- = v_+ = 0$, on a $i = \frac{v_e}{R} = -\frac{v_s}{R_0}$ et $H = -1$. Il s'agit d'un circuit passe-haut.



circuit équivalent en TBF



circuit équivalent en THF

- b) On applique la loi de nœuds exprimée en termes de potentiel à l'entrée inverseuse en notation complexe :

$$\frac{\underline{v_e} - \underline{v_-}}{R + \frac{1}{jC\omega}} = \frac{\underline{v_-} - \underline{v_s}}{R_0}$$

Or comme $v_+ = v_- = 0$, on a finalement :

$$\underline{H} = \frac{\underline{v_s}}{\underline{v_e}} = -\frac{R_0}{R} \frac{1}{1 - j\frac{1}{RC\omega}} = \frac{H_0}{1 - j\frac{\omega_c}{\omega}} = \frac{H_0 j\frac{\omega}{\omega_c}}{1 + j\frac{\omega}{\omega_c}}$$

avec $\omega_c = \frac{1}{RC}$ et $H_0 = -\frac{R_0}{R}$.

2. La fréquence de coupure est :

$$f_c = \frac{\omega_c}{2\pi} = \frac{1}{2\pi RC} = 1,6 \text{ kHz}$$

3. a) La fréquence de coupure à vide est $f_0 = \frac{3 \cdot 10^6}{2 \cdot 10^5} = 15 \text{ Hz}$ soit $\omega_0 = 94 \text{ rad.s}^{-1}$

- b) On applique la loi des nœuds en termes de potentiel à l'entrée inverseuse comme précédemment pour obtenir $\underline{v_-}$. La fonction de transfert s'obtient à partir de la relation de fonction en régime linéaire de l'amplificateur opérationnel : $\underline{v_s} = \underline{\mu} (\underline{v_+} - \underline{v_-})$ avec $\underline{\mu} = \frac{\mu_0}{1 + j\frac{\omega}{\omega_0}}$.

On calcule $\underline{v_-}$ et $\underline{v_+}$:

$$\underline{v_-} = \frac{jR_0 C \omega \underline{v_e} + (1 + jRC\omega) \underline{v_s}}{1 + j(R + R_0) C \omega} \quad \text{et} \quad \underline{v_+} = 0$$

On reporte les expressions de $\underline{v_-}$ et de $\underline{v_+}$ dans la relation $\underline{v_s} = \underline{\mu} (\underline{v_+} - \underline{v_-})$ caractérisant le fonctionnement en régime linéaire de l'amplificateur opérationnel et on trouve :

$$\underline{H'} = \frac{-j\mu_0 R_0 C \omega}{1 + \mu_0 + j \left(\frac{1}{\omega_0} + (R(1 + \mu_0) + R_0) C \right) \omega - \frac{(R + R_0) C}{\omega_0} \omega^2}$$

soit en fonction de $\omega_c = \frac{1}{RC}$ et $H_0 = -\frac{R_0}{R}$:

$$\underline{H'} = \frac{j\mu_0 H_0 \frac{\omega}{\omega_c}}{1 + \mu_0 + j \left(\frac{1}{\omega_0} + \frac{1 - H_0 + \mu_0}{\omega_c} \right) \omega - \frac{1 - H_0}{\omega_0 \omega_c} \omega^2}$$

Alors en tenant compte des ordres de grandeur des constantes à savoir $\mu_0 \gg 1$, $\mu_0 \gg |H_0| = 100$ et $\mu_0 \gg \frac{\omega_c}{\omega_0} = 100$, on obtient finalement :

$$\underline{H'} = \frac{H_0}{1 + j\omega \frac{1 - H_0}{\mu_0 \omega_0} - j\frac{\omega_c}{\omega}}$$

Il s'agit donc d'un filtre passe-bande de fréquence de résonance obtenue par annulation du terme en ω :

$$\omega_r = \sqrt{\frac{\mu_0 \omega_0 \omega_c}{1 - H_0}} = 43,2 \cdot 10^3 \text{ rad.s}^{-1}$$

soit une fréquence de résonance $f_r = 6,9 \text{ kHz}$.

- c) Pour déterminer la bande passante, il faut connaître la facteur de qualité Q du montage. On l'obtient en comparant l'expression de la fonction de transfert obtenue avec la forme canonique

$$\frac{H_0}{1 + jQ \left(\frac{\omega}{\omega_r} - \frac{\omega_r}{\omega} \right)}. \text{ On trouve } Q = \frac{\omega_c}{\omega_r} = 0,23 \text{ et on en déduit la bande passante en fréquence :}$$

$$\Delta f = \frac{f_r}{Q} = 29,7 \text{ kHz}$$

B.3 1. a) Il existe un bouclage entre l'entrée inverseuse et la sortie donc l'amplificateur opérationnel peut fonctionner en régime linéaire.

- b) L'entrée non inverseuse est reliée à la masse et l'amplificateur opérationnel idéal fonctionne en régime linéaire donc le potentiel de l'entrée inverseuse est nul : $v_- = 0$ (l'entrée inverseuse est une masse virtuelle).

On applique la loi des nœuds exprimée en termes de potentiel au point A en utilisant les admittances comme le suggère l'énoncé :

$$\underline{Y}_1 (v_- - v_A) = \underline{Y}_4 v_A + \underline{Y}_3 v_A - \underline{Y}_2 (v_s - v_A)$$

On fait de même à l'entrée inverseuse :

$$\underline{Y}_3 v_A = -\underline{Y}_5 v_s$$

En éliminant le potentiel v_A en A entre les deux expressions, on trouve l'expression demandée :

$$\underline{H} = \frac{v_s}{v_e} = \frac{-\underline{Y}_1 \underline{Y}_3}{\underline{Y}_2 \underline{Y}_3 + \underline{Y}_5 (\underline{Y}_1 + \underline{Y}_2 + \underline{Y}_3 + \underline{Y}_4)}$$

2. a) Pour les résistances, on a $\underline{Y}_i = \frac{1}{R_i}$ et pour les condensateurs $\underline{Y}_i = jC_i \omega$. On peut alors écrire $\underline{H}$ sous la forme :

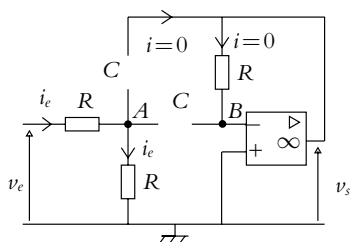
$$\underline{H} = \frac{-\frac{jRC\omega}{2}}{1 + jRC\omega - \frac{R^2 C^2 \omega^2}{2}}$$

Par identification avec l'expression de l'énoncé :

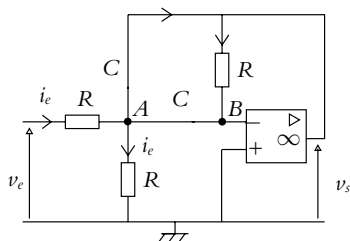
$$\begin{cases} \omega_0 = \frac{\sqrt{2}}{RC} \\ \lambda = \frac{RC\omega_0}{2} = \frac{1}{\sqrt{2}} \\ H_0 = -\frac{1}{2} \end{cases}$$

- b) En TBF, on remplace les condensateurs par des coupe-circuits. Tous les courants sont nuls (sauf i_e) donc $v_s = v_- = 0$ et $\underline{H} = 0$.

Si on fait tendre ω vers 0 dans l'expression de $\underline{H}$, on trouve effectivement 0.



- c) En THF, on remplace les condensateurs par des court-circuits. On a donc $v_A = v_-$ et $v_- = v_+ = 0$ donc $v_A = 0$. De plus, $v_s = v_A$ donc $v_s = 0$ et $\underline{H} = 0$.



Si on fait tendre ω vers ∞ dans l'expression de $\underline{H}$, on obtient :

$$\underline{H} \simeq \frac{2jH_0\lambda\frac{\omega}{\omega_0}}{-\frac{\omega^2}{\omega_0^2}} \simeq 0$$

et on retrouve le même résultat.

Il s'agit donc d'un filtre passe-bande, il ne laisse passer ni les basses fréquences ni les hautes fréquences.

- d) Le module de $\underline{H}$ est :

$$|\underline{H}| = |H_0| \frac{2\lambda\frac{\omega}{\omega_0}}{\left(\left(1 - \frac{\omega^2}{\omega_0^2}\right)^2 + 4\lambda^2\frac{\omega^2}{\omega_0^2}\right)^{1/2}}$$

Quant à l'argument de $\underline{H}$, il vaut :

$$\varphi = \text{Arg}(H_0) + \frac{\pi}{2} - \text{Arg}\left(1 + 2j\lambda\frac{\omega}{\omega_0} - \frac{\omega^2}{\omega_0^2}\right) = \frac{3\pi}{2} - \varphi_D$$

en notant φ_D l'argument du dénominateur.

On a :

$$\tan \varphi_D = \frac{2\lambda\frac{\omega}{\omega_0}}{1 - \frac{\omega^2}{\omega_0^2}}$$

et $\sin \varphi_D > 0$ d'où $\varphi_D \in [0, \pi]$ et $\varphi \in \left[\frac{\pi}{2}, \frac{3\pi}{2}\right]$.

- e) Pour déterminer la pulsation rendant le module de $\underline{H}$ maximal, il vaut mieux écrire H sous la forme :

$$H = \frac{|H_0|}{\left(1 + \frac{\omega_0^2}{4\lambda^2\omega^2} \left(1 - \frac{\omega^2}{\omega_0^2}\right)^2\right)^{1/2}} = \frac{|H_0|}{\left(1 + \left(\frac{\omega_0}{2\lambda\omega} - \frac{\omega}{2\omega_0\lambda}\right)^2\right)^{1/2}}$$

$\underline{H}$ est maximal si le dénominateur est minimal. Ce dernier est la somme de deux nombres positifs dont l'un est constant, il est minimal lorsque la parenthèse est nulle, ce qui se produit pour $\omega = \omega_0$, qui correspond à la pulsation de résonance. On a alors $H_M = |H_0|$.

- f) Les pulsations de coupure à -3 dB correspondent à $H = \frac{H_M}{\sqrt{2}}$ soit :

$$\left(\frac{\omega_0}{\omega} - \frac{\omega}{\omega_0} \right)^2 \frac{1}{(2\lambda)^2} = 1$$

Si on pose $X = \frac{\omega}{\omega_0}$, on obtient deux équations du second degré (en prenant la racine carrée de l'équation ci-dessus) :

$$X^2 \pm 2\lambda X - 1 = 0$$

dont les seules racines physiquement acceptables (car positives) sont :

$$X_1 = -\lambda + \sqrt{1 + \lambda^2} \quad \text{et} \quad X_2 = \lambda + \sqrt{1 + \lambda^2}$$

soit les pulsations

$$\omega_1 = \omega_0 \left(-\lambda + \sqrt{1 + \lambda^2} \right) \quad \text{et} \quad \omega_2 = \omega_0 \left(\lambda + \sqrt{1 + \lambda^2} \right)$$

On en déduit le facteur de qualité :

$$Q = \frac{\omega_2 - \omega_1}{\omega_0} = 2\lambda = \sqrt{2}$$

- g) Pour le tracé du diagramme de Bode, le lecteur est invité à se reporter au paragraphe concernant le filtre passe-bande ; on fera attention pour le tracé de la phase qu'ici H_0 est négatif.
 h) Le circuit se comporte en intégrateur lorsque le diagramme est confondu avec l'asymptote de pente -20 dB/déc soit pour $\omega \gg \omega_0$.
 i) Le circuit se comporte en dérivateur lorsque le diagramme est confondu avec l'asymptote de pente $+20$ dB/déc c'est-à-dire pour $\omega \ll \omega_0$.

B.4 1. a) Pour un temps infini, on obtient le régime permanent qui est ici continu. Par conséquent, les dérivées temporelles sont nulles soit $i = 0$.

b) En écrivant la loi des mailles avec $i = 0$, on a $u_C = E$.

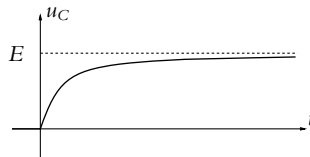
c) Dans ces conditions, la capacité C se comporte comme un interrupteur puisqu'il n'y a pas de courant.

d) En écrivant une loi des mailles ainsi que la relation $i = C \frac{du_C}{dt}$, on obtient $\frac{du_C}{dt} + \frac{u_C}{RC} = \frac{E}{RC}$.

e) La quantité τ est homogène à un temps dit caractéristique donc τ s'exprime en secondes.

f) En résolvant l'équation différentielle avec la condition initiale $u_C = 0$, on obtient $u_C = E \left(1 - \exp \left(-\frac{t}{RC} \right) \right)$ et $u_C = E$ pour un temps infini.

g) On a une asymptote $u = E$ et l'allure est la suivante :



h) L'équation de la tangente à l'origine s'écrit $u_C = \frac{du_C}{dt}(t=0)t + u_C(0) = \frac{E}{RC}t$.

i) L'intersection avec l'asymptote $u_C = E$ s'obtient pour $t = \tau$.

j) Il faut résoudre l'équation $u_C(t) = 0,99E$. On obtient $t_1 = RC \ln 100$.

k) L'intensité s'obtient par la relation $i = C \frac{du_C}{dt} = \frac{E}{R} \exp \left(-\frac{t}{RC} \right)$.

2. a) L'énergie emmagasinée dans C s'écrit $E_C = \int_0^{+\infty} u_C(t)i(t)dt = \frac{CE^2}{2}$.

b) L'énergie dissipée par effet Joule est $E_J = \int_0^{+\infty} Ri^2(t)dt = \frac{CE^2}{2}$.

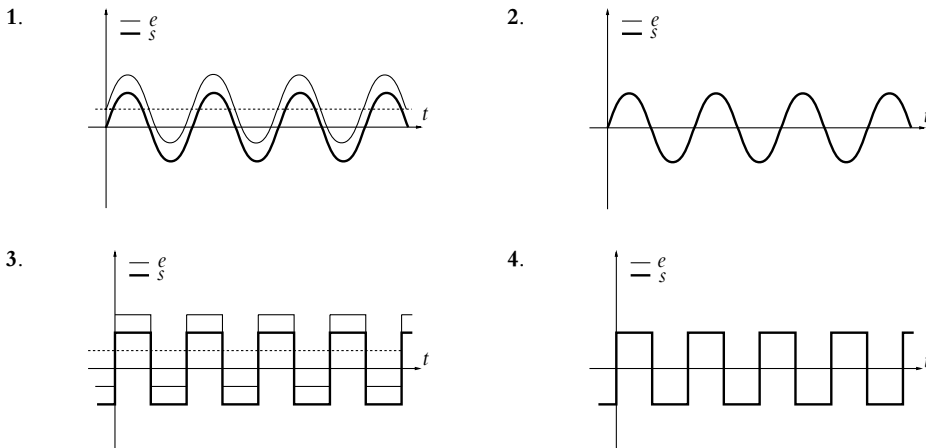
c) L'énergie fournie par le générateur vaut $E_G = \int_0^{+\infty} Ei(t)dt = CE^2$.

- d) On constate que $E_G = E_C + E_J$, ce qui traduit la conservation de l'énergie.
- e) On a $\rho = \frac{CE^2}{2CE^2} = 0,5$.
- f) Par le même raisonnement que précédemment, on obtient $E_{G1} = \frac{CE^2}{4}$ et $E_{C1} = \frac{CE^2}{8}$.
- g) Par la résolution de la même équation différentielle avec une nouvelle condition initiale $u_C = \frac{E}{2}$, on obtient $u_C = E \left(1 - \frac{\exp\left(-\frac{t}{RC}\right)}{2} \right)$.
- h) On en déduit l'intensité $i = C \frac{du_C}{dt} = \frac{E}{2R} \exp\left(-\frac{t}{RC}\right)$.
- i) Pour la seconde phase, on a $E_{G2} = \frac{CE^2}{2}$ et $E_{C2} = \frac{3CE^2}{8}$.
- j) On en déduit $\rho' = 0,66$.
- k) On pourrait faire tendre le rendement énergétique vers 1 en décomposant la charge en N charges de "pas" $\frac{E}{N}$.
3. a) A basses fréquences, la capacité C se comporte comme un interrupteur ouvert donc $s = e$. A hautes fréquences, le comportement est celui d'un fil donc $s = 0$. On en déduit qu'il s'agit d'un filtre passe-bas.
- b) Par la relation du pont diviseur de tension, on en déduit la fonction de transfert $\underline{H} = \frac{1}{1 + jRC\omega} = \frac{H_0}{1 + j\frac{\omega}{\omega_0}}$ avec $H_0 = 1$ et $\omega_0 = \frac{1}{RC}$.
- c) On obtient le gain $G = \frac{1}{\sqrt{1 + R^2 C^2 \omega^2}}$ qui est une fonction décroissante.
- d) Le déphasage s'exprime par $\varphi = -\text{Arctan} RC\omega$ qui est une fonction décroissante.
- e) Lorsque la pulsation ω tend vers 0, le gain G tend vers 1, G_{dB} vers 0 et φ vers 0.
- f) Lorsque la pulsation ω tend vers l'infini, le gain G tend vers 0, G_{dB} vers $-\infty$ et φ vers $-\frac{\pi}{2}$.
- g) On peut écrire l'équation de l'asymptote à l'infini soit
- $$G_{dB} = -20 \log_{10} \sqrt{1 + R^2 C^2 \omega^2} \simeq -20 \log_{10} RC\omega$$
- On a donc une pente à -20 dB par décade.
- h) On a donc le diagramme de Bode d'un filtre passe-bas du premier ordre, on se reportera au cours pour l'allure de ce dernier.
- i) On doit résoudre $G_{dB} \geq -3$ ou $G \geq \frac{G_{max}}{\sqrt{2}} = \frac{1}{\sqrt{2}}$. On obtient $\omega \leq \omega_0$ comme dans le cours.
- j) Lorsque ω tend vers l'infini, on a l'équivalence $\underline{H} \simeq \frac{1}{jRC\omega}$ soit un comportement de type intégrateur.
4. a) On a $\underline{\epsilon} = \underline{v_+} - \underline{v_-} = 0$ car l'amplificateur opérationnel est supposé idéal et en régime linéaire. On a $\underline{v_+} = 0$ et par le théorème de Millman, on obtient $\underline{v_-} = \frac{\frac{e}{R} + \frac{s}{Z}}{\frac{1}{R} + \frac{1}{Z}}$ avec $\underline{Z} = \frac{R'}{1 + jR'C\omega}$ soit
- $$\underline{H} = \frac{-R'}{R(1 + jR'C\omega)}.$$
- b) On en déduit l'équation différentielle $\frac{ds}{dt} + \frac{s}{R'C} = -\frac{e}{RC}$.
- c) On a un comportement intégrateur si $R'C\omega \gg 1$ et $s = -\frac{E}{RC\omega} \sin \omega t$.
- d) On a l'approximation $\underline{H} \simeq -\frac{R'}{R}$ si $R'C\omega \ll 1$ et le gain vaut alors $-\frac{R'}{R}$.

- e) Le potentiel v_- peut être considéré comme une masse virtuelle puisque $v_+ = 0$ et qu'on a $\epsilon = 0$.
- f) La résolution de l'équation différentielle avec la condition initiale $s(0) = 0$ (on l'obtient avec $u_C = s$ en utilisant la question précédente) donne $s = \frac{R'E}{R} \left(\exp\left(-\frac{t}{R'C}\right) - 1 \right)$.
- g) En effectuant un développement limité, on obtient $s \simeq -\frac{E}{RC}t$ et le comportement est de type intégrateur.
- h) Ce sera le cas si $t \ll R'C$.
- i) L'amplificateur n'est pas en régime saturé donc $|s| \leq V_{\text{sat}}$ et on doit donc avoir $E \leq \frac{R}{R'} V_{\text{sat}}$.

Chapitre 19

A.1 En supposant que le filtre est idéal, il ne laisse passer que les fréquences supérieures à 100 Hz et coupe toutes les autres. Pour des signaux de 2 kHz, cela revient à conserver le fondamental et toutes les harmoniques et à supprimer la composante continue. On obtient donc les signaux d'entrée centrés autour de la valeur 0 V dont les allures sont les suivantes :

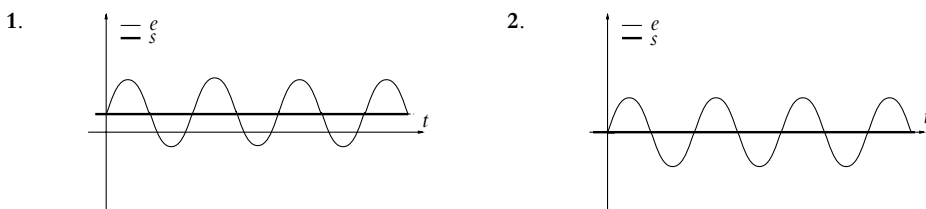


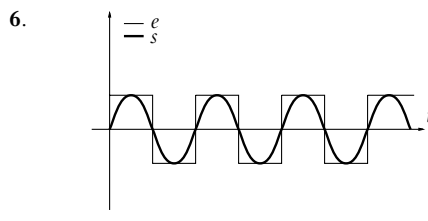
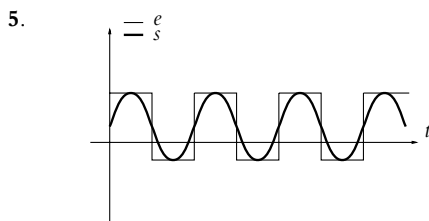
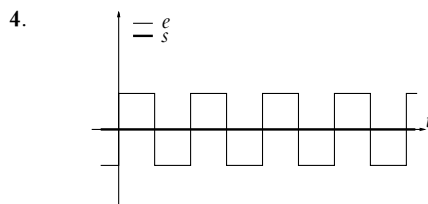
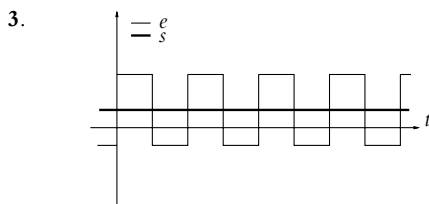
A.2 En supposant que le filtre est idéal, il ne laisse passer que les fréquences inférieures à 100 Hz et coupe toutes les autres.

Pour des signaux de 2 kHz, cela revient à ne conserver que la composante continue et à supprimer le fondamental et toutes les harmoniques. On obtient un signal continu.

Pour des signaux de 75 Hz, cela revient à ne conserver que la composante continue et le fondamental et à supprimer toutes les autres harmoniques : on obtient donc une sinusoïde, pas nécessairement centrée sur 0.

Les allures sont donc les suivantes :





B.1 1. L'impédance équivalente à R_1 , L_1 et C_1 en parallèle s'écrit :

$$\underline{Z} = \frac{1}{\frac{1}{R_1} + \frac{1}{jL_1\omega} + jC_1\omega} = \frac{jR_1L_1\omega}{R_1 - R_1L_1C_1\omega^2 + jL_1\omega}$$

et la tension

$$\underline{v_s} = \underline{Z}I_0 = \frac{jR_1L_1\omega}{R_1 - R_1L_1C_1\omega^2 + jL_1\omega} \underline{v_e}$$

On en déduit la fonction de transfert :

$$\begin{aligned} \underline{H} &= \frac{jR_1L_1\omega}{R_1 - R_1L_1C_1\omega^2 + jL_1\omega} = \frac{R_1s}{1 + jR_1\sqrt{\frac{C_1}{L_1}}\left(\omega\sqrt{L_1C_1} - \frac{1}{\omega\sqrt{L_1C_1}}\right)} \\ &= \frac{A}{1 + jQ\left(x - \frac{1}{x}\right)} \end{aligned}$$

avec $A = sR_1$, $Q = R_1\sqrt{\frac{C_1}{L_1}}$ et $\omega_0 = \frac{1}{\sqrt{L_1C_1}}$.

2. Ce filtre est un filtre passe-bande car quand $\omega \rightarrow 0$ ou quand $\omega \rightarrow \infty$, $\underline{H} \rightarrow 0$ et quand $\omega \rightarrow \omega_0$, $\underline{H} \rightarrow A$.

3. La bande passante est la zone de pulsations telle que :

$$\frac{A}{\sqrt{1 + Q^2\left(x - \frac{1}{x}\right)^2}} \geq \frac{A}{\sqrt{2}}$$

soit

$$Q^2\left(x - \frac{1}{x}\right)^2 \leq 1$$

dont la résolution (le calcul se mène de la même manière qu'au paragraphe sur la résonance aux bornes de R d'un circuit R, L, C série) donne :

$$\frac{-1 + \sqrt{1 + 4Q^2}}{2Q} \leq x \leq \frac{1 + \sqrt{1 + 4Q^2}}{2Q}$$

et une bande passante :

$$\Delta\omega = \frac{\omega_0}{Q}$$

$$4. \quad \delta_n = \frac{\tau_{n,s}}{\tau_{n,e}} = \frac{V_{n,s} V_{1,e}}{V_{1,s} V_{n,e}} = \frac{1}{\sqrt{1 + Q^2 \left(n - \frac{1}{n}\right)^2}}$$

en appliquant les définitions de l'énoncé. On notera que $x = n$ avec les notations proposées.

5. A partir de l'expression de δ_n , on en déduit :

$$Q = \frac{\sqrt{1 - \delta_n^2}}{\delta_n \left(n - \frac{1}{n}\right)}$$

Il suffit alors de faire l'application numérique pour $n = 2$ et $\delta_n = 0,1$: $Q = 6,63$ puis pour $\delta_n = 0,01$: $Q = 21,1$.

6. Comme on a choisi un facteur de qualité qui permet de réduire d'un facteur 100 l'amplitude de l'harmonique d'ordre 2, le filtre ne garde que le fondamental dont la pulsation est la pulsation de résonance du filtre. On obtient donc une sinusoïde de même fréquence que le créneau.

- B.2** 1. On remplace dans l'expression générale obtenue à l'exercice B.3 du chapitre précédent les admittances $\underline{Z}_1$, $\underline{Z}_2$ et $\underline{Z}_4$ par $\frac{1}{R}$ et les admittances $\underline{Z}_3$ par $jC_1\omega$ et $\underline{Z}_5$ par $jC_2\omega$. On obtient :

$$\underline{G} = -\frac{1}{1 + 3jRC_1\omega - R^2C_1C_2\omega^2}$$

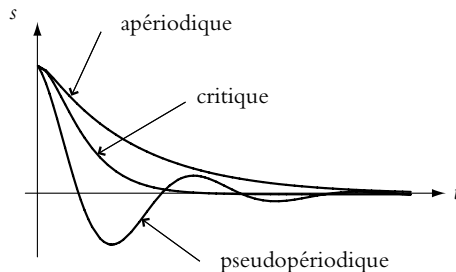
Par identification avec l'expression proposée dans l'énoncé, on trouve : $\omega_0 = \frac{1}{R\sqrt{C_1C_2}}$,

$$m = \frac{3}{2}\sqrt{\frac{C_1}{C_2}} \text{ et } G_0 = -1.$$

2. On utilise le passage de la fonction de transfert à l'équation différentielle grâce à l'équivalence entre la multiplication par $j\omega$ en notation complexe et la dérivation par rapport au temps en notation réelle. On en déduit :

$$\frac{1}{\omega_0^2} \frac{d^2s}{dt^2} + \frac{2m}{\omega_0} \frac{ds}{dt} + s = G_0 e$$

3. La solution générale de l'équation différentielle homogène associée décrit le régime libre. L'équation caractéristique associée $\frac{1}{\omega_0^2}r^2 + \frac{2m}{\omega_0}r + 1 = 0$ admet pour discriminant réduit $\Delta' = \frac{m^2 - 1}{\omega_0^2}$. On a donc trois régimes possibles pour le régime libre : apériodique si $\Delta' > 0$ ou $m > 1$, critique si $\Delta' = 0$ ou $m = 1$ et pseudopériodique si $\Delta' < 0$ ou $m < 1$. Les allures sont les suivantes :



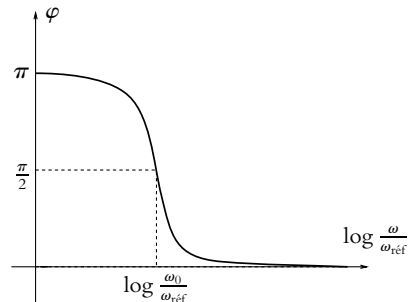
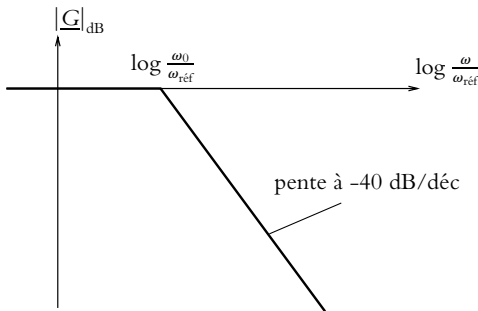
4. La tension aux bornes d'un condensateur est une fonction continue du temps. Par conséquent, la tension aux bornes de C_1 qui est égale à la tension de sortie du montage ($\underline{v}_B = 0$) est continue à $t = 0$. On en déduit $s(0^+) = 0$ puisque les condensateurs sont initialement déchargés. $t \rightarrow \infty$ correspond au régime permanent qui sera décrit par une solution particulière constante de l'équation différentielle. On obtient : $s(+\infty) = G_0 E_0$.

5. On rappelle que $m = \frac{3}{2}\sqrt{\frac{C_1}{C_2}}$. Si on veut $m = \frac{\sqrt{2}}{2} < 1$, il faut donc choisir C_2 tel que $C_2 = \frac{9}{2}C_1$. Le régime libre est alors apériodique d'après ce qui a été fait auparavant.

6. On souhaite $10^3 \text{ Hz} \leq f_0 = \frac{\sqrt{2}}{6\pi RC_1} \leq 4 \cdot 10^3 \text{ Hz}$, ce qui impose $\frac{\sqrt{2}}{24\pi 10^3 C_1} \leq R \leq \frac{\sqrt{2}}{6\pi 10^3 C_1}$.

7. Pour obtenir les variations du gain, on étudie la fonction $f(\omega) = \left(1 - \frac{\omega^2}{\omega_0^2}\right)^2 + 4m^2 \frac{\omega^2}{\omega_0^2}$ qui est croissante (on peut par exemple établir que sa dérivée est positive). On en déduit que $|G|$ est décroissante et qu'il n'y a pas de phénomène de résonance. L'étude des limites donne les équations des asymptotes suivantes :

- si $\omega \rightarrow 0$: $G_{\text{dB}} = 0$,
 - si $\omega \rightarrow +\infty$: $G_{\text{dB}} = 40 \log \omega_0 - 40 \log \omega$.
- Leur intersection est obtenue pour $\omega = \omega_0$.



Pour l'étude de la phase : $\varphi = \text{Arg } \underline{G} = \pi - \varphi_D$ en utilisant les notations de l'exercice B.3 du chapitre précédent. Or on peut établir que φ_D est décroissante (par exemple en étudiant sa dérivée par rapport au temps) : φ est décroissante de π à 0.

8. On a un comportement intégrateur lorsque la fonction de transfert peut s'identifier à la forme : $\frac{1}{j\frac{\omega}{\omega_0}}$,

ce qui sera le cas ici au voisinage de ω_0 .

Pour avoir un comportement dérivateur, il faudrait que la fonction de transfert puisse s'identifier à la forme $j\frac{\omega}{\omega_0}$, ce qui n'est jamais le cas ici.

9. a) Comme $1,5 < 2$, la fréquence du signal est dans la bande passante donc $s = -e$ c'est-à-dire un déphasage de π et une amplitude identique.
- b) Pour le fondamental, on a les mêmes résultats qu'à la question précédente. Pour la première harmonique obtenue pour $n = 3$, son amplitude vaut $0,11a$ et le gain à cette fréquence est également de $0,11$: l'amplitude en sortie du filtre est donc de 1% par rapport à celle du fondamental. On peut la négliger et le signal en sortie est une sinusoïde en opposition de phase par rapport au créneau.
- c) On effectue le même raisonnement et on obtient le même résultat qu'à la question précédente.
10. A l'ordre 2, la pente de l'asymptote obtenue pour $\omega \rightarrow +\infty$ est de -40 dB/déc : on a donc une diminution plus rapide des composantes de hautes fréquences qu'on souhaite supprimer.

B.3 1. L'amplificateur opérationnel étant idéal et en régime linéaire, on a $\epsilon = v_+ - v_- = 0$. Comme la borne non inverseuse est reliée à la masse, on en déduit que le potentiel à l'entrée inverseuse est nul. On applique le théorème de Millman en A puis en $v_- = 0$ d'après ce qui précède. On obtient $\frac{v_-}{Y_3 + Y_5} = \frac{Y_3 v_A + Y_5 \underline{s}}{Y_1 \underline{\epsilon} + Y_4 \underline{s}}$ et $\frac{v_A}{Y_1 + Y_2 + Y_3 + Y_4} = \frac{Y_1 \underline{\epsilon} + Y_4 \underline{s}}{Y_1 + Y_2 + Y_3 + Y_4}$. En éliminant $\frac{v_A}{Y_1 + Y_2 + Y_3 + Y_4}$ entre ces deux expressions, on obtient l'expression demandée en posant $\underline{S} = \frac{Y_1}{Y_1 + Y_2 + Y_3 + Y_4}$.

2. a) Avec $\frac{Y_1}{-1} = \frac{Y_3}{-1} = \frac{Y_4}{-1} = \frac{1}{R}$, $\frac{Y_2}{-1} = jC_2\omega$ et $\frac{Y_5}{-1} = jC_1\omega$, on obtient $\underline{H} = \frac{-1}{1 + 3jRC_1\omega - R^2C_1C_2\omega^2}$.

b) Il s'agit donc d'un filtre passe-bas du second ordre.

c) On peut mettre la fonction de transfert sous sa forme canonique à savoir $\frac{H_0}{1 + j\frac{1}{Q}\frac{\omega}{\omega_0} - \left(\frac{\omega}{\omega_0}\right)^2}$

avec $H_0 = -1$, $\omega_0 = \frac{1}{R\sqrt{C_1C_2}}$ et $Q = \frac{1}{3}\sqrt{\frac{C_2}{C_1}}$.

d) On en déduit les valeurs de C_1 et C_2 à partir des expressions de ω_0 et Q soit $C_2 = \frac{3Q}{R\omega_0}$ et $C_1 = \frac{1}{3RQ\omega_0}$.

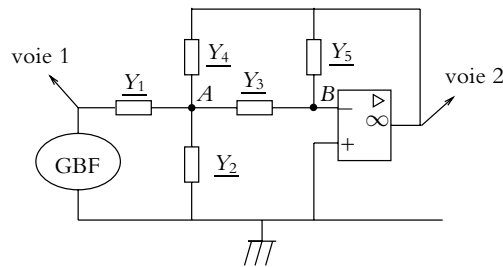
e) A partir de l'expression de ω_0 , on en déduit l'incertitude relative

$$\frac{\Delta\omega_0}{\omega_0} = \frac{\Delta R}{R} + \frac{\Delta C}{C} = \frac{\Delta C}{C}$$

du fait de la précision infinie sur R et de la précision sur C à 5 %.

f) L'application numérique donne $C_2 = 16,9 \text{ nF}$ et $C_1 = 3,75 \text{ nF}$.

3. a) Le branchement expérimental est le suivant :

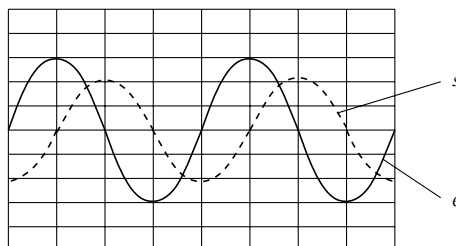


b) Pour réaliser le tracé du diagramme de Bode, on visualise les deux voies en fonction du temps, on mesure l'amplitude (en tenant compte du calibre et en mesurant le nombre de carreaux ou en utilisant les outils intégrés) et du déphasage (par la méthode des neuf carreaux ou de Lissajous ou en utilisant les outils intégrés).

c) Une période correspond à 4 carreaux, chaque carreau représentant $12,5 \mu\text{s}$. On en déduit que la période est $T = 50 \mu\text{s}$ et la fréquence $f = \frac{1}{T} = 20 \text{ kHz}$. Quant à l'amplitude d'entrée, elle est de trois carreaux correspondant chacun à $1,0 \text{ V}$ donc l'amplitude est de $3,0 \text{ V}$.

d) Pour une fréquence $f = f_0$, la fonction de transfert s'écrit $\underline{H} = jQ$. En faisant l'application numérique, on en déduit la valeur du gain $0,707$ et du déphasage $\frac{\pi}{2}$.

e) Par conséquent, l'amplitude de sortie vaut $2,12 \text{ V}$ et on obtient l'allure suivante :



f) Pour observer le diagramme de Bode avec ce dispositif à l'aide d'un crêteau, il faut prendre un signal de très faible fréquence afin que les harmoniques intéressantes pour le filtre (au voisinage de 20 kHz) soient à peu près constantes d'après le rappel de l'énoncé. On envoie ce signal à l'entrée du filtre et on fait l'analyse de Fourier de la sortie : le spectre de la sortie pour les harmoniques élevées donne l'allure du gain.

B.4 1. En appliquant la relation du pont diviseur de tension, on obtient $\underline{H} = \frac{jRC\omega}{1 + jRC\omega}$.

2. Il s'agit d'un filtre passe-haut du premier ordre.

3. Le gain s'écrit $G = |\underline{H}| = \frac{1}{\sqrt{1 + \frac{1}{(RC\omega)^2}}}$. La dérivée de $f(\omega) = 1 + \frac{1}{R^2C^2\omega^2}$ est

$f'(\omega) = -\frac{2}{R^2C^2\omega^3} < 0$. On en déduit que f est décroissante et que le gain G est une fonction croissante.

4. Les asymptotes sont $G_{dB} \simeq 20 \log_{10} RC\omega$ pour ω tendant vers 0 et $G_{dB} \simeq 0$ pour ω infini. Ces deux équations définissent le diagramme asymptotique.

5. Pour déterminer la bande passante, il faut résoudre $\frac{1}{\sqrt{1 + \frac{1}{R^2C^2\omega^2}}} \geq \frac{1}{\sqrt{2}}$. On obtient

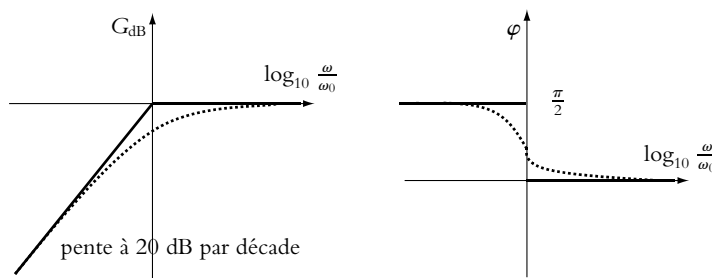
$$\omega \geq \omega_0 = \frac{1}{RC}.$$

6. Le déphasage vaut $\varphi = -\text{Arg}\left(1 - j\frac{1}{RC\omega}\right) = -\varphi'$ avec $\tan \varphi' = \frac{1}{RC\omega}$ et $\cos \varphi'$ du signe de 1 soit positif. Par conséquent, on a $\varphi' \in \left[-\frac{\pi}{2}, \frac{\pi}{2}\right]$. On en déduit $\varphi = \text{Arctan} \frac{1}{RC\omega}$.

7. Soit $f(\omega) = \frac{1}{RC\omega}$. On a $f'(\omega) = -\frac{1}{RC\omega^2} < 0$ donc la fonction f est décroissante et $\varphi = \text{Arctanf}(\omega)$ l'est aussi.

8. Pour les limites, on a $\lim_{\omega \rightarrow 0} \varphi = \frac{\pi}{2}$ et $\lim_{\omega \rightarrow +\infty} \varphi = 0$.

9. Les diagrammes de Bode réel et asymptotique sont représentés respectivement en pointillés et en trait plein :



10. Pour obtenir la valeur de C cherchée, on doit effectuer la résolution de $-20 \log_{10} \sqrt{1 + \frac{1}{R^2C^2\omega^2}} = -6$.

On obtient $C = \frac{1}{2\pi Rf\sqrt{10^{0,6} - 1}} = 9,2 \cdot 10^{-9}$ F.

11. Pour le montage 2, on applique à nouveau la relation du pont diviseur de tension et on obtient

$\underline{H} = \frac{\underline{Z}_{\parallel}}{\underline{Z}_s + \underline{Z}_{\parallel}}$ avec $\underline{Z}_s = \frac{1 + jRC\omega}{jC\omega}$ et $\underline{Z}_{\parallel} = \frac{R}{1 + jRC\omega}$. En explicitant les valeurs des impédances

dans la fonction de transfert, on en déduit $\underline{H} = \frac{jRC\omega}{1 + 3jRC\omega - R^2C^2\omega^2}$.

12. Il s'agit d'un filtre passe-bande du second ordre.

13. La fonction de transfert peut s'écrire sous la forme $\underline{H} = \frac{A}{1 + jQ \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)}$ avec $A = Q = \frac{1}{3}$ et

$$\omega_0 = \frac{1}{RC}.$$

14. Pour obtenir les variations du gain, on étudie $f(\omega) = 1 + Q^2 \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)^2$. Cette fonction vérifie $f(\omega) \geq 1$ pour toute valeur de ω et le cas d'égalité est obtenu pour $\omega = \omega_0$. On a donc un phénomène de résonance pour cette pulsation.

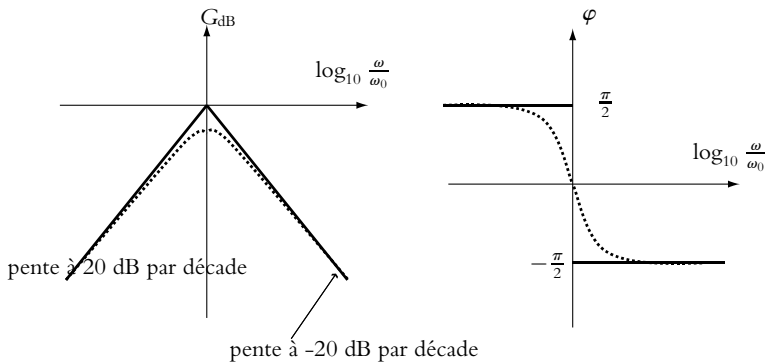
15. Pour le diagramme de Bode asymptotique, on a $G_{dB} \simeq 20 \log_{10} \frac{A}{Q} + 20 \log_{10} \frac{\omega}{\omega_0}$ pour $\omega \rightarrow 0$ et $G_{dB} \simeq 20 \log_{10} \frac{A}{Q} - 20 \log_{10} \frac{\omega}{\omega_0}$ pour $\omega \rightarrow \infty$. Ces asymptotes ont une intersection pour $\omega = \omega_0$.

16. Le déphasage est $\psi = \text{Arg} A - \text{Arg} \left(1 + jQ \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right) \right) = -\psi'$ avec $\tan \psi' = Q \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)$ et $\cos \psi'$ du signe de 1 soit positif.
On en déduit $\psi = -\text{Arctan} Q \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)$.

17. Les variations du déphasage sont obtenues par l'étude de $f(\omega) = -Q \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)$ dont la dérivée est $f'(\omega) = Q \left(\frac{1}{\omega_0} + \frac{\omega_0}{\omega^2} \right) > 0$. On en déduit que f et ψ sont décroissantes.

18. Les valeurs limites du déphasage sont $\lim_{\omega \rightarrow 0} \psi = \frac{\pi}{2}$ et $\lim_{\omega \rightarrow +\infty} \psi = -\frac{\pi}{2}$.

19. Les diagrammes de Bode réel et asymptotique sont représentés respectivement en pointillés et en trait plein :



20. L'application numérique donne $f_0 = \frac{1}{2\pi RC} = 1,7 \text{ kHz}$.

21. Pour une fréquence $f = 1,0 \text{ kHz}$, on a un gain de $G = 0,30$ pour le fondamental (fréquence f), $G = 0,33$ pour l'harmonique d'ordre 2, $G = 0,30$ pour celle d'ordre 3, $G = 0,28$ pour celle d'ordre 4, $G = 0,25$ pour celle d'ordre 5, $G = 0,22$ pour celle d'ordre 6, $G = 0,20$ pour celle d'ordre 7, $G = 0,18$ pour celle d'ordre 8, $G = 0,17$ pour celle d'ordre 9 et $G = 0,15$ pour celle d'ordre 10.

22. Le signal de sortie est quasi-identique à celui d'entrée avec une amplitude divisée approximativement par 3 puisque le gain est sensiblement égal pour toutes les premières harmoniques.

23. Pour déterminer la bande passante, on résout $\frac{A}{\sqrt{1 + Q^2 \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)^2}} \geq \frac{A}{\sqrt{2}}$ soit

$$(Q\omega^2 - \omega_0\omega - Q\omega_0^2)(Q\omega^2 + \omega_0\omega - Q\omega_0^2) \leq 0.$$

On obtient $\omega_0 \frac{-1 + \sqrt{1 + 4Q^2}}{2Q} \leq \omega \leq \omega_0 \frac{1 + \sqrt{1 + 4Q^2}}{2Q}$. On en déduit la largeur de la bande passante $\Delta\omega = \frac{\omega_0}{Q} = \frac{3}{RC}$.

24. Pour un circuit R, L, C série, on a $\underline{i} = \frac{\underline{\varepsilon}}{R + j \left(L\omega - \frac{1}{C\omega} \right)}$. La résonance en intensité se produit

pour la fréquence $f_0 = \frac{1}{2\pi\sqrt{LC}} = 1,0 \text{ MHz}$.

25. Pour le circuit R, L, C série, l'impédance est $\underline{Z}_s = R + j \left(L\omega - \frac{1}{C\omega} \right)$. Pour $\omega = \omega_0$, on a $\underline{Z}_s = R = 1,0 \cdot 10^4 \Omega$.

Pour le circuit proposé, on a l'impédance

$$\begin{aligned} \underline{Z}_{||} &= \frac{1}{jC\omega + \frac{1}{R + jL\omega}} = \frac{R + jL\omega}{1 - LC\omega^2 + jRC\omega} = \frac{L}{RC} - j\sqrt{\frac{L}{C}} \\ &= 250 \cdot 10^6 - 1,58 \cdot 10^6 \Omega \end{aligned}$$

soit $\underline{Z}_{||} \simeq \frac{L}{RC} = 250 \cdot 10^6 \Omega$ pour $\omega = \omega_0$.

26. On applique la relation du pont diviseur de tension et $\underline{\varepsilon} = \frac{\underline{Z}_{||}}{\underline{Z}_s + \underline{Z}_{||}} \underline{\varepsilon}$ soit compte tenu des valeurs numériques un gain de quasi 1 : on a une amplitude de 0,5 V et un déphasage quasinel.

Chapitre 20

A.1 1. Un ampèremètre se branche en série dans la branche où on veut mesurer le courant. Pour ne pas perturber la mesure, il faut que le courant dans la branche soit imposé par les éléments de celle-ci et non par l'ampèremètre. Cela implique que la résistance de l'appareil doit être faible devant celle des éléments de la branche.

2. Un voltmètre se branche en parallèle aux bornes du dipôle dont on veut mesurer la tension. Pour ne pas perturber la mesure, il faut que cette tension soit imposée par le dipôle et non par le voltmètre. Cela implique que la résistance de l'appareil doit être grande devant celle du dipôle.

A.2 1. Le montage 2 est inintéressant : on ne mesure pas la tension aux bornes de la résistance. Les deux autres montages en revanche sont utilisables.

2. Pour le montage 1 : $u_{\text{mes}} = (R + R_A) i_{\text{mes}}$ donc $R_{\text{mes}} = \frac{u_{\text{mes}}}{i_{\text{mes}}} = R + R_A$.

Pour le montage 2 : $u_{\text{mes}} = R_V i_V = R(i_{\text{mes}} - i_V)$ en notant i_V l'intensité du courant traversant le voltmètre. On en déduit : $R_{\text{mes}} = \frac{u_{\text{mes}}}{i_{\text{mes}}} = R_V \frac{i_V}{i_{\text{mes}}}$. Or $(R_V + R) i_V = R i_{\text{mes}}$ donc : $R_{\text{mes}} = \frac{RR_V}{R + R_V}$.

3. Lorsque que la valeur de la résistance R est faible devant celle de R_A , le montage 1 donne R_A comme valeur de résistance, ce qui est évidemment une valeur erronée. Dans ce montage, on surestime la valeur de la tension alors qu'on a la bonne valeur pour l'intensité. Il faut donc que la chute de tension aux bornes de l'ampèremètre soit négligeable devant celle aux bornes de la résistance donc que $R \gg R_A$.

A l'inverse, lorsque la valeur de la résistance R est grande devant celle de R_V , le montage 2 donne R_V comme valeur de résistance, ce qui est évidemment faux. Dans ce montage, on surestime la valeur de l'intensité alors qu'on a la bonne valeur pour la tension. Il faut donc que l'intensité dans le voltmètre soit négligeable devant celle dans la résistance donc que $R \ll R_V$.

On utilisera donc plutôt le montage 1 pour les fortes valeurs de résistance et le 2 pour les faibles valeurs.

- B.1** 1. Lorsqu'on ouvre l'interrupteur, la capacité C se trouve uniquement relié à la résistance R_f dans laquelle elle se décharge. Ceci explique la diminution de la tension aux bornes du condensateur.
2. La décharge de C dans R_f vérifie l'équation différentielle suivante :

$$\frac{du}{dt} + \frac{u}{R_f C} = 0$$

en notant u la tension aux bornes du condensateur. Compte tenu des conditions initiales : $u(0) = E$, la solution s'écrit : $u = E \exp\left(-\frac{t}{R_f C}\right)$. Or à $t = T$, $u = E'$. On en déduit donc :

$$R_f = \frac{T}{C} \ln \frac{E}{E'}$$

- B.2** 1. On observe une charge du condensateur : une loi des mailles permet d'obtenir l'équation différentielle : $RC \frac{du_C}{dt} + u_C = E$ dont la solution est, compte tenu des conditions initiales,

$$u_C = E \left(1 - \exp\left(-\frac{t}{RC}\right)\right) = E \left(1 - \exp\left(-\frac{t}{\tau}\right)\right)$$

La tension du condensateur tend vers E .

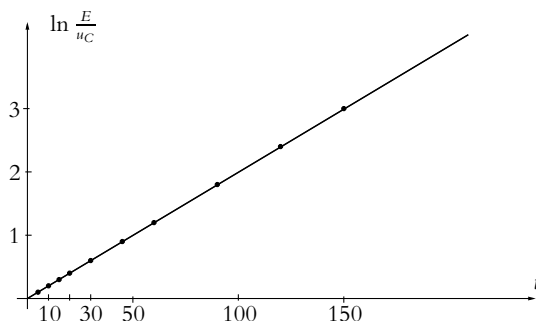
2. L'équation différentielle s'écrit maintenant : $RC \frac{du_C}{dt} + u_C = 0$. La solution est : $u_C = E \exp\left(-\frac{t}{\tau}\right)$. On en déduit :

$$\ln \frac{E}{u_C} = \frac{t}{\tau}$$

On obtient donc une droite de pente $\frac{1}{\tau}$.

t (s)	5	10	15	20	30	45	60	90	120	150
u_C (V)	9,05	8,19	7,41	6,70	5,49	4,07	3,01	1,65	0,91	0,50
$\ln \frac{E}{u_C}$	0,10	0,20	0,30	0,40	0,60	0,90	1,20	1,80	2,40	3,00

3.

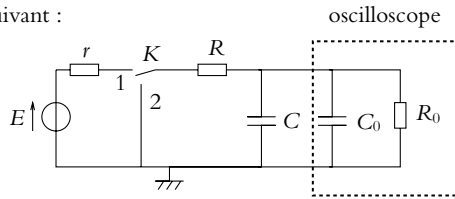


On obtient bien une droite comme l'annonce la théorie mais la pente vaut 0,02 au lieu de 0,01. L'écart est significatif.

4. On n'a pas pris en compte les paramètres du voltmètre. Si on en tient compte, le nouveau montage est obtenu en plaçant en parallèle avec C la résistance R_V du voltmètre. On peut retrouver le montage initial en remplaçant R par $R_{eq} = \frac{RR_V}{R + R_V}$ et la pente p vaut alors $\frac{1}{R_{eq}C}$. On peut en déduire la valeur

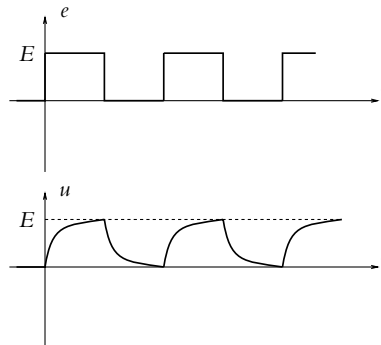
de $R_V = \frac{R}{pRC - 1} = 10 \text{ M}\Omega$. Ceci est tout à fait possible compte tenu des ordres de grandeur usuels des résistances des voltmètres.

5. L'effet de R_V aurait été négligeable si $R \ll R_V$ alors $R_{eq} \simeq R$ et on retrouve les résultats théoriques.
6. Un GBF peut être modélisé par une source idéale de tension e en série avec une résistance interne r de 50Ω .
7. Pour l'oscilloscope, on distingue le mode DC où l'impédance d'entrée est constituée d'un condensateur de capacité C_0 de l'ordre de 50 pF en parallèle avec résistance R_0 de $1 \text{ M}\Omega$ du mode AC où l'impédance d'entrée comporte en plus un condensateur de grande capacité en série avec l'impédance d'entrée du mode DC. Dans la suite de l'exercice, on utilise l'oscilloscope en mode DC.
8. On a donc le schéma suivant :

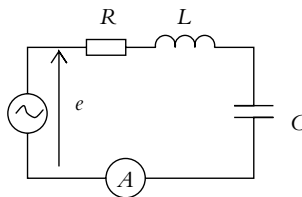


En procédant par équivalence entre le modèle de Thévenin et le modèle de Norton d'un dipôle linéaire, on montre qu'il suffit de remplacer E par $\frac{R_0 E}{R + r + R_0}$, R par $\frac{R_0 (r + R)}{r + R + R_0}$ et C par $C + C_0$ pour obtenir un schéma identique à celui du circuit étudié au début de l'exercice.

9. On vérifie avec les données numériques que $r = 50 \Omega \ll R = 10 \text{ k}\Omega \ll R_0 = 1 \text{ M}\Omega$ et $C_0 = 50 \text{ pF} \ll C = 10 \text{ nF}$. On retrouve dans ce cas $R' = R$, $C' = C$ et $E' = E$.
10. On prend sur la sortie 50Ω un signal créneau entre 0 V et $E \text{ V}$.
11. Pour observer complètement la décharge et la recharge, il faut que $T \gg \tau$ soit $f \ll 10 \text{ kHz}$. On observe alors :



B.3 1.



2. On utilise la notation complexe et :

$$i = \frac{e}{Z} = \frac{e}{R + j \left(L\omega - \frac{1}{C\omega} \right)}$$

On en déduit le module : $I = \frac{E}{\sqrt{R^2 + \left(L\omega - \frac{1}{C\omega}\right)^2}}$ et le déphasage φ tel que

$$\tan \varphi = \frac{-L\omega + \frac{1}{C\omega}}{R} \text{ et } \varphi \in \left[-\frac{\pi}{2}, \frac{\pi}{2}\right] \text{ car } \cos \varphi > 0.$$

3. On considère ici I comme une fonction de C : I est maximale si $f(C) = R^2 + \left(L\omega - \frac{1}{C\omega}\right)^2$ est minimale. Or $f(C) > R^2$ pour toute valeur de C et $f(C) = R^2$ pour $C = C_0 = \frac{1}{L\omega^2}$. La fonction $I(C)$ présente bien un maximum.

4. Pour obtenir les valeurs de C_1 et de C_2 , il faut résoudre :

$$I = \frac{E}{\sqrt{R^2 + \left(L\omega - \frac{1}{C\omega}\right)^2}} = \frac{I_0}{\sqrt{2}} \quad \text{avec} \quad I_0 = \frac{E}{R}.$$

Cela conduit à l'équation :

$$R^2 + \left(L\omega - \frac{1}{C\omega}\right)^2 = 2R^2 \quad \text{soit} \quad L\omega - \frac{1}{C\omega} = \pm R$$

On en déduit :

$$C_1 = \frac{1}{L\omega^2 + R\omega} \quad \text{et} \quad C_2 = \frac{1}{L\omega^2 - R\omega}$$

5. C_1 et C_2 sont supposés proches, ce qui sera possible si $L\omega^2 \ll R\omega$ soit $L\omega \ll R$.

6. Dans le cadre de cette approximation : $C_1 \simeq \frac{1}{L\omega^2} \left(1 - \frac{R}{L\omega}\right)$ et $C_2 \simeq \frac{1}{L\omega^2} \left(1 + \frac{R}{L\omega}\right)$ donc :

$$C_0 - C_1 = C_2 - C_0 = \frac{R}{L^2\omega^3}$$

7. Lorsque $C = C_0$ alors $\tan \varphi = 0$ donc comme $\cos \varphi > 0$, on en déduit : $\varphi = 0$: il n'y a pas de déphasage.

8. Comme $\frac{1}{C_1\omega} = L\omega + R$ et $\frac{1}{C_2\omega} = L\omega - R$, on en déduit : $\tan \varphi_1 = 1$ et $\tan \varphi_2 = -1$ soit $\varphi_1 = \frac{\pi}{4}$ et $\varphi_2 = -\frac{\pi}{4}$ puisque $\cos \varphi > 0$. On a donc des valeurs quadrantales.

9. Pour obtenir C_0 , on cherche la valeur maximale de I à l'ampèremètre ou à l'oscilloscope. Comme on a établi qu'alors on a un déphasage nul, il est plus facile avec un oscilloscope d'appliquer la méthode de Lissajous en ajoutant une faible résistance au circuit aux bornes de laquelle on obtient l'image de l'intensité (la faible valeur de la résistance est liée à la nécessité de ne pas perturber le reste du circuit).

10. La mesure de C_0 est délicate car il faut soit ajouter une petite résistance soit chercher un maximum, ce qui conduit à des incertitudes importantes. D'autre part, pour déterminer R et L , il faut deux mesures, ce que donne la détermination de C_1 et de C_2 .

11. Le coefficient de surtension aux bornes de C vaut :

$$Q = \frac{u_C}{E} = \frac{E}{RC_0\omega E} = \frac{L\omega E}{RE} = \frac{u_L}{E}$$

qui est donc également le coefficient de surtension aux bornes de L .

12. Comme $Q = \frac{L\omega}{R}$ et que $C_2 - C_1 = \frac{2RC_0}{L\omega} = \frac{2C_0}{Q}$, on en déduit : $Q = \frac{C_2 + C_1}{C_2 - C_1}$ puisque

$$C_0 = \frac{C_1 + C_2}{2}.$$

13. De $C_0 = \frac{C_1 + C_2}{2} = \frac{1}{L\omega^2}$, on tire : $L = \frac{1}{2\pi^2(C_1 + C_2)f^2}$. Alors en utilisant l'expression de Q , on en déduit : $R = \frac{C_2 - C_1}{\pi f(C_1 + C_2)^2}$

14. Si on tient compte de la résistance de l'ampèremètre, R devient $R + R_A$. Par conséquent, la valeur de L ne change pas et $R = \frac{C_2 - C_1}{\pi f(C_1 + C_2)^2} - R_A$.
15. On obtient : $Q = \frac{C_1 + C_2}{C_1 - C_2} = 25$, $R = \frac{C_2 - C_1}{\pi f(C_1 + C_2)^2} = 0,51 \, \Omega$ en négligeant la résistance de l'ampèremètre, $R = \frac{C_2 - C_1}{\pi f(C_1 + C_2)^2} - R_A = 0,41 \, \Omega$ en tenant compte et $L = \frac{1}{2\pi^2(C_1 + C_2)f^2} = 20 \, \mu\text{H}$.
16. L'approximation $R = 0,5 \, \Omega \ll L\omega = 12,6 \, \Omega$ est justifiée. En revanche, R est du même ordre de grandeur que R_A et l'approximation $R_A \ll R$ ne peut être faite.

Chapitre 21

- A.1** 1. On sait que le signal est obtenu en sortie d'un amplificateur opérationnel. On peut imaginer qu'il s'agit de la sortie de l'un des montages suivants :
- montage suiveur avec un créneau en entrée,
 - montage suiveur avec un triangle en entrée,
 - montage intégrateur avec un créneau en entrée,
 - montage amplificateur (inverseur ou non) avec un créneau en entrée,
 - montage amplificateur (inverseur ou non) avec un triangle en entrée.
2. La différence entre les deux observations est un écrêtage non symétrique dans le deuxième cas qui est sans doute lié à la présence d'une tension de décalage du signal d'entrée ou d'une tension d'offset de l'amplificateur opérationnel (plus probable avec un montage amplificateur). Pour les différentes solutions proposées, on met en évidence les défauts suivants :
- on a un retard à la transition d'une valeur à l'autre lié au slew rate et éventuellement un phénomène de saturation de l'amplificateur opérationnel,
 - si le signal d'entrée est un triangle, on observe alors le phénomène de saturation de l'amplificateur opérationnel,
 - pour l'intégrateur, le défaut est toujours un phénomène de saturation,
 - comme pour le suiveur, on met en évidence le slew rate et éventuellement un phénomène de saturation d'autant plus probable que l'amplification est forte,
 - ici seul un phénomène de saturation peut expliquer ce qui est observé.

- A.2** 1. Si on néglige la tension d'offset, on obtient : $v_+ = 0$ et en régime linéaire pour un amplificateur idéal, $\epsilon = v_+ - v_- = 0$ donc $v_- = 0$ et finalement $s = 0$.
2. En tenant compte de l'offset, on a toujours $v_+ = 0$ mais

$$\epsilon' = \epsilon + V_d = V_d = v_+ - v_- = -v_-.$$

Or l'application de la loi des nœuds exprimée en termes de potentiel à l'entrée inverseuse v_- donne :

$$\frac{0 - v_-}{R_1} + \frac{s - v_-}{R_2} = 0 \quad \text{soit} \quad v_- = \frac{R_1 s}{R_1 + R_2}$$

On en déduit que :

$$s = - \left(1 + \frac{R_2}{R_1} \right) V_d$$

3. En inversant cette relation, on trouve :

$$V_d = \frac{R_1 s}{R_1 + R_2}$$

La mesure de la tension de sortie permet donc d'accéder à la valeur de la tension de décalage ou d'offset.

► **Remarque** : pour le 741, le constructeur garantit une tension d'offset maximale entre 5 et 7,5 mV suivant la température et pour le TL081 entre 6 et 9 mV. Pour pouvoir effectuer la mesure, il faut que la tension de sortie ne dépasse pas la tension de saturation $V_{\text{sat}} \simeq 13 \text{ V}$. Cela impose que le choix des résistances R_1 et R_2 vérifie :

$$\frac{R_1}{R_1 + R_2} = \frac{10 - -3}{13} \simeq 10^{-3}$$

Par exemple, on pourrait choisir $R_1 = 100 \Omega$ et $R_2 = 100 \text{ k}\Omega$. On notera que ce choix évite de prendre une résistance trop grande qui risque de poser des problèmes vis-à-vis de la résistance d'entrée de l'oscilloscope.

A.3 1. Si on suppose $I_{0+} = I_{0-}$ (ce qui n'est pas forcément le cas),

$$I_{0-} = \frac{dq}{dt} \quad \text{donc} \quad q = I_{0-}t = Cu_c = -Cs$$

en supposant qu'initialement la capacité n'est pas chargée et qu'à $t = 0$ on ferme l'interrupteur K . On en déduit donc

$$I_{0-} = \frac{-CU}{\Delta t}$$

en notant U la valeur de la tension de sortie obtenue après un temps Δt . La capacité se charge en effet du fait de l'existence du courant de polarisation et la mesure du temps nécessaire pour atteindre une tension de sortie U qu'on se fixe *a priori* permet d'en déduire la valeur des courants de polarisation.

► **Remarque** : Cette mesure est assez délicate à mettre en œuvre si on ne dispose pas du mode ROLL (Cf. chapitre sur l'instrumentation électrique) d'un oscilloscope numérique. On répétera alors l'expérience plusieurs fois dans les mêmes conditions et on prendra la valeur moyenne des résultats obtenus.

2. Une méthode pour compenser ces courants de polarisation pourrait être d'utiliser un générateur de courant externe réglable qui "annulerait" les courants I_{0+} et I_{0-} .
3. Pour le montage proposé, on a :

$$\begin{cases} v_+ = RI_{0+} \\ v_- = \frac{I_{0-} + \frac{e}{R_1} + \frac{s}{R_2}}{\frac{1}{R_1} + \frac{1}{R_2}} = \frac{R_2 e + R_1 s + R_1 R_2 I_{0-}}{R_1 + R_2} \end{cases}$$

L'amplificateur opérationnel voit donc les résistances R_+ et R_- suivantes aux entrées non inverseuse et inverseuse :

$$R_+ = R \quad \text{et} \quad R_- = \frac{R_1 R_2}{R_1 + R_2}$$

quand e est reliée à la masse et qu'on a $s = 0$ comme le veut alors le modèle idéal de l'amplificateur opérationnel. On choisit donc sur ce montage d'insérer un résistor de résistance

$$R = \frac{R_1 R_2}{R_1 + R_2}$$

entre l'entrée inverseuse et la masse.

A.4 On formule l'hypothèse que les amplificateurs opérationnels fonctionnent en régime linéaire sinon la notion d'impédance n'a pas de sens. Comme les amplificateurs opérationnels sont supposés idéaux, on a $\epsilon = 0$ dans le cas du régime linéaire. Pour l'amplificateur opérationnel AO2, en utilisant le théorème de Millman pour calculer le potentiel à l'entrée inverseuse, on obtient $\underline{s}_2 = -\frac{1}{jR_4 C \omega} \underline{s}_1$. Pour l'amplificateur opérationnel AO1, le potentiel à l'entrée non inverseuse est égal à $\underline{e}$. D'autre part, par la relation du pont

diviseur de tension, on a $\underline{e} = -\frac{R_5}{R_5 + R_3} \underline{s}_1$. La loi des nœuds sur l'entrée inverseuse s'écrit

$$i_e = i_{R_1} + i_{R_2} = \frac{\underline{e} - \underline{s}_1}{R_1} + \frac{\underline{e} - \underline{s}_2}{R_2} = \left(\frac{2}{R_1} + \frac{R_3}{R_1 R_5} + \frac{1}{R_2} + \frac{R_3 + R_5}{j R_2 R_4 R_5 C \omega} \right) \underline{e}$$

On a donc comme impédance d'entrée l'association en parallèle d'une résistance $R_e = \frac{R_1 R_2 R_5}{2 R_2 R_5 + R_2 R_3 + R_1 R_5}$ et d'une inductance $L_e = \frac{R_2 R_4 R_5 C}{R_3 + R_5}$.

A.5 1. On formule l'hypothèse que l'amplificateur opérationnel fonctionne en régime linéaire sinon la notion de fonction de transfert n'a pas de sens. Comme l'amplificateur est idéal, on en déduit que $\epsilon = 0$ soit $\underline{v}_+ = \underline{s}$. Par la relation du pont diviseur de tension, on a $\underline{v}_+ = \frac{\underline{v}_A}{1 + j R C \omega}$ en notant A le point commun à R , R' et C' . En appliquant le théorème de Millman en A , on obtient $\underline{v}_A = \frac{R \underline{e} + R' \underline{s} + j R R' C' \omega \underline{s}}{R + R' + j R R' C' \omega}$. Il suffit alors d'éliminer $\underline{v}_A$ entre les deux relations qui viennent d'être obtenues pour en déduire la fonction de transfert $\underline{H} = \frac{1}{1 + j(R + R') C \omega - R R' C C' \omega^2}$.

2. On calcule son module $|\underline{H}| = \frac{1}{\sqrt{(1 - R R' C C' \omega^2)^2 + (R + R')^2 C^2 \omega^2}}$ soit en développant $|\underline{H}| = \frac{1}{\sqrt{1 + ((R + R')^2 C^2 - 2 R R' C C') \omega^2 + R^2 R'^2 C^2 C'^2 \omega^4}}$. Pour obtenir la forme souhaitée, il

faut annuler le terme en ω^2 c'est-à-dire avoir $(R + R')^2 C = 2 R R' C'$ et $\omega_0 = \frac{1}{\sqrt{R R' C C'}}$.

3. L'étude de $f(\omega) = 1 + \left(\frac{\omega}{\omega_0}\right)^4$ montre qu'il s'agit d'une fonction croissante : le gain est donc une fonction décroissante de ω . Les asymptotes ont pour équation $G_{dB} \simeq 0$ pour ω tendant vers 0 et $G_{dB} \simeq -40 \log_{10} \frac{\omega}{\omega_0}$ pour ω infini (on a alors une pente à -40 dB par décade).

4. La bande passante est définie par les valeurs de ω telles que $G \geq \frac{G_{max}}{\sqrt{2}}$. La résolution de cette inéquation donne $\omega \leq \omega_0$.

5. Pour déterminer le déphasage, il convient de bien revenir à la fonction de transfert dans laquelle on insère éventuellement la condition obtenue sur les valeurs des composants. On a $\tan \varphi = -\frac{(R + R') C \omega}{1 - R R' C C' \omega^2}$ et $\cos \varphi > 0$ pour $\omega \leq \omega_0$ et < 0 sinon. On en déduit $\varphi = -\text{Arctan} \frac{(R + R') C \omega}{1 - R R' C C' \omega^2}$ si $\omega \leq \omega_0$ et $\varphi = \pi - \text{Arctan} \frac{(R + R') C \omega}{1 - R R' C C' \omega^2}$ sinon. L'étude de ces fonctions montre que le déphasage φ est une fonction décroissante de 0 à $-\frac{\pi}{2}$ puis de $\frac{3\pi}{2}$ à π .

A.6 1. Un amplificateur opérationnel est idéal si les courants de polarisation sont nuls (ou de façon équivalente si l'impédance d'entrée est infinie), si la résistance de sortie est nulle et si le gain différentiel statique est infini.

2. Un amplificateur opérationnel fonctionne en régime linéaire si on a la relation $s = \mu_0 \epsilon$ avec $\epsilon = v_+ - v_-$.

3. En appliquant le théorème de Millman au second amplificateur opérationnel avec $\epsilon = 0$ (il est idéal et en régime linéaire), on a $v_- = 0 = \frac{u + u'}{2}$. On en déduit la tension de sortie du second amplificateur opérationnel $u' = -u$.

On fait la même chose pour le premier amplificateur opérationnel et on obtient $v_- = e = \frac{R_3 u + R_1 s}{R_1 + R_3}$.

On est en court-circuit si $s = 0$. Or le courant de sortie s'exprime par

$$i_s = i_{R_2} + i_{R_3} = \frac{-(R_1 + R_3)e + R_1 s}{R_2 R_3} - \frac{s}{R_2} + \frac{e - s}{R_3}$$

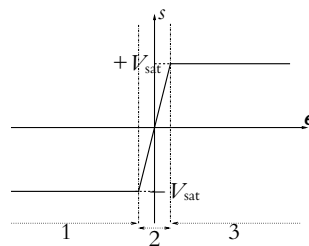
On en déduit le courant de court-circuit $i_{s,CC} = \frac{R_2 - R_1 - R_3}{R_2 R_3} e$.

4. En circuit ouvert, on a $i_s = 0$ et exceptionnellement on peut appliquer le théorème de Millman en sortie de l'amplificateur opérationnel soit $s = \frac{R_2 e + R_3 u'}{R_2 + R_3}$. On en déduit $s_0 = \frac{R_2 - R_1 - R_3}{R_2 + R_3 - R_1} e$.
5. Vu de la sortie, le montage est équivalent à une source de courant de courant dont le courant de court-circuit est $i_{s,CC}$ et la résistance interne $R_N = \frac{s_0}{i_{s,CC}} = \frac{R_2 R_3}{R_2 + R_3 - R_1}$. On a donc un générateur de Norton.
6. On utilise les relations d'association en série et en parallèle des dipôles pour en déduire l'impédance $\underline{Z} = R_4 + \frac{R_5}{1 + jR_5 C \omega}$.

7. La norme de l'impédance s'écrit $|\underline{Z}| = \sqrt{\frac{(R_4 + R_5)^2 + (R_4 R_5 C \omega)^2}{1 + (R_5 C \omega)^2}} = 7,8 \text{ k}\Omega$.

8. La tangente du déphasage $\tan \varphi$ est le rapport de la partie imaginaire par la partie réelle soit $\tan \varphi = -\frac{R_5^2 C \omega}{R_4 + R_5 + R_4 R_5^2 C^2 \omega^2}$.
9. La quantité $\cos \varphi$ est du signe de la partie réelle donc ici positif. On en déduit $\varphi = -\text{Arctan} \frac{R_5^2 C \omega}{R_4 + R_5 + R_4 R_5^2 C^2 \omega^2} = -39,8^\circ = -0,695 \text{ rad}$.
10. L'intensité efficace vaut $I = \frac{U}{Z} = 12,8 \text{ mA}$ et la puissance $P_{AB} = UI \cos \varphi = 1,28 \text{ W}$.
11. La puissance dissipée dans la résistance R_4 vaut $P_{R_4} = R_4 I^2 = 0,164 \text{ W}$.
12. De même, la puissance consommée dans la capacité est nulle soit $P_C = 0$.

B.1 1. La caractéristique statique de l'amplificateur opérationnel est la suivante :



2. Un amplificateur opérationnel est idéal si les courants de polarisation sont nuls (ou de façon équivalente si l'impédance d'entrée est infinie), si la résistance de sortie est nulle et si le gain différentiel statique est infini.
3. Un amplificateur opérationnel fonctionne en régime linéaire si on a la relation $s = \mu_0 \epsilon$ avec $\epsilon = v_+ - v_-$.
4. La sortie de l'amplificateur opérationnel ne prend pas de valeur infinie, ce qui implique avec l'hypothèse que μ_0 est infini que ϵ doit être nul.
5. Par la relation du pont diviseur de tension, on a $\underline{v_-} = \frac{P}{P+Q} \underline{s}$. Or $\underline{v_+} = \underline{v_-} = \underline{\epsilon}$. On en déduit donc

$$\frac{\underline{s}}{\underline{\epsilon}} = 1 + \frac{Q}{P}.$$

6. L'amplificateur opérationnel étant idéal, le courant d'entrée $i_e = i_+$ est nul.
7. Compte tenu des réponses précédentes, on a R_e infini, $R_s = 0$ du fait du caractère idéal de l'amplificateur opérationnel et $G = 1 + \frac{Q}{P}$.
8. On a $G = 1 + \frac{Q}{P} = 1$ si $P \gg Q$ soit par exemple pour $P = 100 \text{ k}\Omega$ et $Q = 1 \text{ k}\Omega$, ces valeurs ont été choisies pour que ni la résistance interne du GBF (50Ω) ni l'impédance d'entrée de l'oscilloscope ($1 \text{ M}\Omega$) ne modifie le comportement du montage.
- Par ailleurs, on a $G' = 1 + \frac{Q'}{P'} \simeq 2$ si $P' = Q'$ soit par exemple $P' = Q' = 10 \text{ k}\Omega$ avec les mêmes contraintes que précédemment pour le choix des valeurs.

9. On applique le théorème de Millman à l'entrée non inverseuse du premier amplificateur opé-

$$\text{rationnel } \underline{v}_{+1} = \frac{\frac{\underline{\epsilon}}{R} + jC\omega \underline{s}}{\frac{1}{R} + jC\omega}$$

$$\underline{v}_{-1} = \frac{P}{Q + P} \underline{u} = \frac{\underline{u}}{G}.$$

De même, avec la relation du pont diviseur de tension aux deux entrées du deuxième amplificateur opérationnel, on obtient $\underline{v}_{+2} = \frac{\underline{u}}{1 + jRC\omega}$ et $\underline{v}_{-2} = \frac{\underline{s}}{G'}$.

En égalant les entrées des deux amplificateurs opérationnels puisque $\underline{\epsilon} = 0$ (ils sont idéaux et fonctionnent en régime linéaire) et en éliminant $\underline{u}$ entre les deux égalités obtenues, on en déduit $\frac{\underline{\epsilon} + jRC\omega \underline{s}}{1 + jRC\omega} = \frac{1 + jRC\omega}{GG'} \underline{s}$ et la fonction de transfert du montage

$$\underline{H} = \frac{GG'}{1 + jRC\omega (2 - GG') - (RC\omega)^2}.$$

10. Avec les quantités que l'énoncé demande d'introduire, on obtient

$$H = \frac{GG'}{\sqrt{(1 - x^2)^2 + x^2 (2 - GG')^2}}$$

11. L'étude de $f(x) = (1 - x^2)^2 + x^2 (2 - GG')^2$ permet d'établir que :

- si $GG' \notin [2 - \sqrt{2}, 2 + \sqrt{2}]$, la dérivée $f'(x) > 0$ et le gain est une fonction décroissante de ω ,
- si $GG' \in [2 - \sqrt{2}, 2 + \sqrt{2}]$, la dérivée $f'(x) > 0$ pour $x \geq x_0 = \sqrt{\frac{2 - (2 - GG')^2}{2}}$ et le gain est une fonction croissante puis décroissante. Le maximum du gain vaut

$$H_{\max} = \frac{GG'}{|2 - GG'| \sqrt{1 - \frac{(2 - GG')^2}{4}}}.$$

12. On pose $GG' = 2 + \epsilon$ et on effectue un développement limité en ϵ . On obtient

$$H_{\max} \simeq \frac{2 + \epsilon}{\epsilon \sqrt{1 - \frac{\epsilon^2}{4}}} = \frac{GG'}{2 - GG'}.$$

13. Pour obtenir la bande passante, il faut résoudre

$$\frac{GG'}{\sqrt{(1 - x^2)^2 + x^2 (2 - GG')^2}} \geq \frac{GG'}{\sqrt{2} (2 - GG')}$$

Si $GG' \in \left[2 - \frac{1}{\sqrt{2}}, 2 + \frac{1}{\sqrt{2}}\right]$, on a deux solutions

$$x_1 = \sqrt{\frac{2 - (2 - GG')^2 - \sqrt{(2 - (2 - GG')^2)^2 - 4 + 8(GG' - 2)^2}}{2}}$$

et

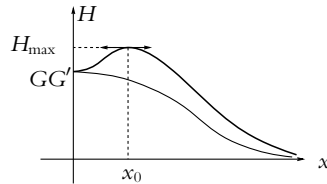
$$x_2 = \sqrt{\frac{2 - (2 - GG')^2 + \sqrt{(2 - (2 - GG')^2)^2 - 4 + 8(GG' - 2)^2}}{2}}$$

On en déduit la largeur de la bande passante $\Delta\omega = \omega_0(x_2 - x_1)$.

Si $GG' \notin \left[2 - \frac{1}{\sqrt{2}}, 2 + \frac{1}{\sqrt{2}}\right]$, il n'y a qu'une seule solution possible x_2 et la largeur de la bande passante est alors $\Delta\omega = \omega_0 x_2$.

14. On effectue un développement limité en ϵ comme précédemment et on en déduit $A = 2 - GG'$.

15. Les différentes allures du gain sont donc :



16. On peut faire varier A en même temps que G' c'est-à-dire en modifiant les valeurs de P' et Q' .

17. Les applications numériques donnent $\omega_0 = 10,0 \text{ rad.s}^{-1}$, $G = 1,00$, $G' = 1,80$ et $GG' = 1,80$. Pour la valeur maximale du gain, on a $H_{\max} = 9,00$ en utilisant l'approximation et $9,05$ sans : on peut donc utiliser les approximations. Enfin le rapport x_0 vaut $x_0 = 0,99$.

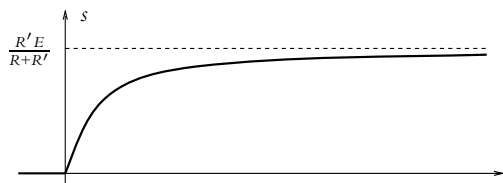
18. On est dans le cas où $\omega \gg \omega_0 \simeq \omega_r$ donc le fondamental et les harmoniques sont coupés : il ne reste que la composante continue ou valeur moyenne soit $\langle e \rangle = aE_0 = 0,20 \text{ V}$. On a donc $\langle s \rangle = H(0) \langle e \rangle = GG' \langle e \rangle = 0,36 \text{ V}$.

B.2 1. a) En appliquant les lois de Kirchhoff (lois des nœuds et des mailles) ou en utilisant les complexes, on obtient l'équation différentielle $RR'C \frac{ds}{dt} + (R + R')s = R'e$.

b) En régime continu, les dérivées temporelles sont nulles donc $i_C = C \frac{ds}{dt} = 0$. Alors en utilisant la relation du pont diviseur de tension, on en déduit $s = \frac{R'e}{R + R'} = 0$ pour $t < 0$ (l'alimentation n'est alors pas branchée). En appliquant la continuité de la tension aux bornes d'une capacité, on a $s(0) = 0$.

c) La solution est la somme d'une solution particulière constante $\frac{R'E}{R + R'}$ et de la solution générale $S \exp\left(-\frac{R + R'}{RR'C}t\right)$. La constante S est déterminée grâce à la condition initiale $s(0) = 0$ soit finalement $s = \frac{R'E}{R + R'} \left(1 - \exp\left(-\frac{R + R'}{RR'C}t\right)\right)$.

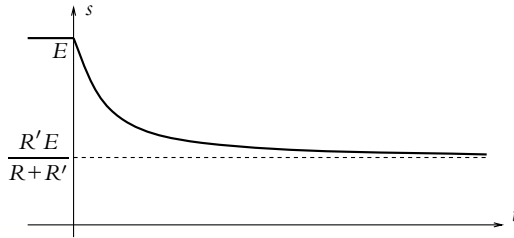
d) L'allure de l'évolution temporelle de s est :



- e) En écrivant une loi des mailles, on en déduit

$$u = e - s = \frac{RE + R'E \exp\left(-\frac{R + R'}{RR'C}t\right)}{R + R'}.$$

- f) L'allure de la tension u est :



2. a) En traduisant l'équation différentielle de la partie précédente en notation complexe, on obtient la

fonction de transfert $\underline{H} = \frac{R'}{R + R' + jRR'C\omega}$.

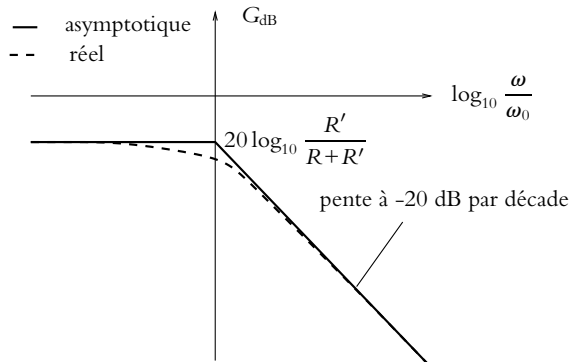
- b) Il s'agit d'un filtre passe-bas du premier ordre.

- c) On peut écrire la fonction de transfert sous sa forme canonique $\underline{H} = \frac{H_0}{1 + j\frac{\omega}{\omega_0}}$ avec $f_0 = \frac{R + R'}{2\pi RR'C}$

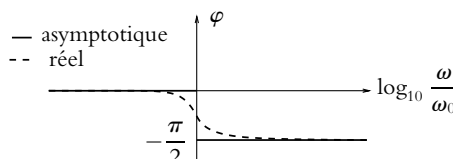
et $H_0 = \frac{R'}{R + R'}$.

- d) On montre comme dans le cours que le gain est décroissant et que le diagramme asymptotique admet pour équations $G_{dB} = 20\log_{10} \frac{R'}{R + R'}$ pour ω tendant vers 0 et

$G_{dB} = 20\log_{10} \frac{R'}{R + R'} - 20\log_{10} \frac{\omega}{\omega_0}$ pour ω infini. On a donc la courbe suivante :



- e) Le déphasage est $\varphi = -\text{Arg}\left(1 + j\frac{\omega}{\omega_0}\right) = -\text{Arctan}\frac{\omega}{\omega_0}$. On montre qu'il s'agit d'une fonction décroissante de 0 à $-\frac{\pi}{2}$. L'allure obtenue est :



3. a) Un amplificateur opérationnel est idéal si les courants de polarisation sont nuls (ou de façon équivalente si l'impédance d'entrée est infinie), si la résistance de sortie est nulle et si le gain différentiel statique est infini.
- b) Un amplificateur opérationnel fonctionne en régime linéaire si on a la relation $s = \mu_0 \epsilon$ avec $\epsilon = v_+ - v_-$.
- c) La sortie de l'amplificateur opérationnel ne prend pas de valeur infinie, ce qui implique avec l'hypothèse que μ_0 est infini que ϵ doit être nul.
- d) On applique le théorème de Millman à l'entrée inverseuse de l'amplificateur opérationnel soit

$$\underline{v_-} = \frac{\frac{\underline{\epsilon}}{R} + \frac{\underline{s}}{Z_{||}}}{\frac{1}{R} + \frac{1}{Z_{||}}} \text{ avec } Z_{||} = \frac{R}{1 + jRC\omega}. \text{ Comme } \underline{v_+} = 0, \text{ on en déduit la fonction de transfert}$$

$$\underline{H} = -\frac{1}{1 + jRC\omega}.$$
- e) On a le même comportement que précédemment avec $H_0 = -1$ et $\omega_0 = \frac{1}{RC}$. On a donc une fréquence de coupure plus basse, un gain plus grand et surtout une modification du déphasage $\varphi = \pi - \text{Arctan} \frac{\omega}{\omega_0}$.
- f) L'intérêt du montage à amplificateur opérationnel est l'indépendance du comportement par rapport à la charge R' .

4. a) On calcule les différentes intégrales et on obtient $e(t) = -\frac{4E}{\pi} \sum_{p=0}^{+\infty} \frac{\sin(2p+1)\omega t}{2p+1}$.
- b) Chaque terme est multiplié par $|\underline{H}((2p+1)\omega)|$ et on ajoute un déphasage. On en déduit l'expression $s(t) = \frac{4E}{\pi} \sum_{p=0}^{+\infty} \frac{\sin((2p+1)\omega t - \text{Arctan}((2p+1)RC\omega))}{(2p+1)\sqrt{1 + ((2p+1)RC\omega)^2}}$.
- c) Il est impossible d'avoir $\underline{H} \simeq jK\omega$ donc on n'a pas de comportement dérivateur possible.
- d) Pour $\omega \gg \omega_0$, on a $\underline{H} \simeq \frac{1}{jK\omega}$ et un comportement intégrateur.
- e) On en déduit la condition $RC \gg \frac{T}{2\pi}$ pour avoir un tel comportement.

B.3 1. L'intensité s'écrit $\underline{i} = \frac{\underline{u}}{\underline{Z}}$ donc l'amplitude demandée est $I_M = \frac{V_M}{\sqrt{R^2 + X^2}}$ et le déphasage $\varphi = \text{Arg} \underline{Z} = \text{Arctan} \frac{X}{R}$.
Si $\underline{Z}$ a un comportement inductif, le déphasage est positif.

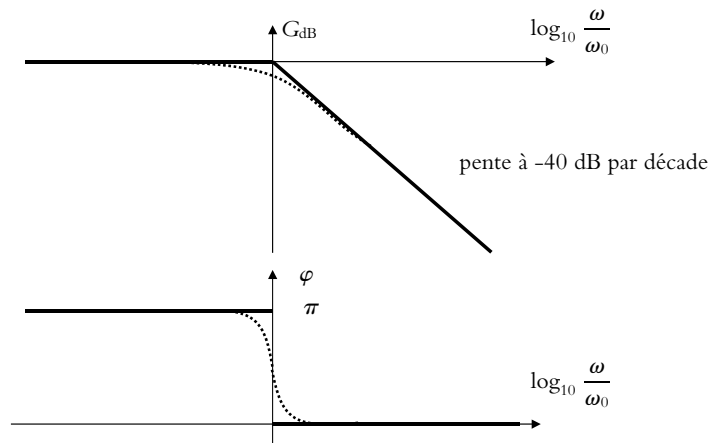
2. La puissance moyenne est définie par $P = \langle p(t) \rangle = \langle u(t)i(t) \rangle = \frac{V_M I_M}{2} \cos \varphi$.
3. Les applications numériques donnent $I_M = 28,6 \text{ A}$, $\varphi = 63,4^\circ$ et $P = 2,05 \text{ kW}$.
4. La valeur moyenne de $\langle v_2(t) \rangle$ peut se mettre sous la forme

$$\langle v_2(t) \rangle = -\langle k v_a(t) v_b(t) \rangle = -K k_a k_b P$$

On obtient bien la forme demandée avec $k_m = -K k_a k_b$.

5. L'application numérique donne $P = 3,75 \text{ kW}$.
6. On applique le théorème de Millman aux points A et B, ce qui donne après simplification $v_A = \frac{v_1 + v_2 + v_B}{3 + jRC_2\omega}$ et $v_B = \frac{v_A + jRC_1\omega v_2}{1 + kRC_1\omega}$. Comme l'amplificateur opérationnel est idéal et qu'il fonctionne en régime linéaire, on a $\epsilon = 0$ soit $v_B = 0$. Finalement on en déduit la fonction de transfert $\underline{H} = \frac{-1}{1 + 3jRC_1\omega - R^2 C_1 C_2 \omega^2}$.

7. On obtient la forme canonique proposée pour $H_0 = -1$, $m = \frac{3}{2} \sqrt{\frac{C_1}{C_2}}$ et $\omega_0 = \frac{1}{R\sqrt{C_1 C_2}}$.
8. Les valeurs des capacités doivent vérifier $C_1 C_2 = \frac{1}{E\pi^2 R^2 f_0^2}$ et $\frac{C_1}{C_2} = \frac{2}{9}$. On en déduit $C_1 = \frac{2}{9} C_2$ et $C_2 = \frac{3}{2\pi\sqrt{2}Rf_0} = 144 \text{ nF}$. Finalement $C_1 = 31,9 \text{ nF}$.
9. Le gain s'écrit $G = \frac{1}{\sqrt{(1 - R^2 C_1 C_2 \omega^2)^2 + 9R^2 C_1^2 \omega^2}}$.
L'étude de $f(\omega) = (1 - R^2 C_1 C_2 \omega^2)^2 + 9R^2 C_1^2 \omega^2$ de dérivée $f'(\omega) = 4R^4 C_1^2 C_2 \omega^3 > 0$ conduit à la conclusion que le gain G est une fonction décroissante de ω .
10. Asymptotiquement $G_{\text{dB}} = 0$ pour ω tendant vers 0 et $G_{\text{dB}} = -40 \log_{10} R\sqrt{C_1 C_2} \omega$ de pente -40 dB par décade pour ω infini. L'intersection des asymptotes se produit pour $\omega = \omega_0$ avec $G_{\text{dB}}(\omega_0) = -10 \log_{10} 2 = -3,0 \text{ dB}$.
11. Le déphasage s'écrit $\varphi = \text{Arg}(-1) - \text{Arg}(1 + 3jRC_1\omega - R^2 C_1 C_2 \omega^2)$ soit $\varphi = -\text{Arctan} \frac{3RC_1\omega}{1 - R^2 C_1 C_2 \omega^2}$ si $\omega > \omega_0$ et $\varphi = \pi - \text{Arctan} \frac{3RC_1\omega}{1 - R^2 C_1 C_2 \omega^2}$ si $\omega < \omega_0$ (on calcule la tangente du déphasage et on étudie le signe de son cosinus pour obtenir son expression).
12. On étudie $f(\omega) = \frac{3RC_1\omega}{1 - R^2 C_1 C_2 \omega^2}$ de dérivée $f'(\omega) = \frac{3RC_1(1 + R^2 C_1 C_2 \omega^2)}{(1 - R^2 C_1 C_2 \omega^2)^2} > 0$. On en déduit que le déphasage φ est une fonction décroissante.
13. Les valeurs particulières sont $\varphi(0) = \pi$, $\varphi(\omega_0) = \frac{\pi}{2}$ et $\varphi(\infty) = 0$.
14. On en déduit le diagramme de Bode suivant :



15. En utilisant le lien entre fonction de transfert et équation différentielle, on déduit de ce qui précède l'équation différentielle $\frac{1}{\omega_0^2} \frac{d^2 v_2}{dt^2} + \frac{2m}{\omega_0} \frac{dv_2}{dt} + v_2 = H_0 v_1$.
16. La tension v_1 vaut $v_1(t) = K v_a(t) v_b(t) = \frac{K k_a k_b V_M I_M}{2} (\cos(2\omega t - \varphi) + \cos \varphi)$. On a donc :
- i. une composante continue $\frac{K k_a k_b V_M I_M \cos \varphi}{2}$ qui donne en sortie

$$H_0 \frac{K k_a k_b V_M I_M \cos \varphi}{2} = - \langle v_1(t) \rangle = -1,38 \text{ V}$$

- ii. une composante sinusoïdale $\frac{Kk_a k_b V_M I_M}{2} \cos(2\omega t - \varphi)$ qui donne en sortie
- $$G(2\omega) \frac{Kk_a k_b V_M I_M}{2} \cos(2\omega t - \varphi)$$

L'amplitude de la composante continue est négligeable devant la composante continu. On peut noter qu'on commet une erreur de 0,23 %.

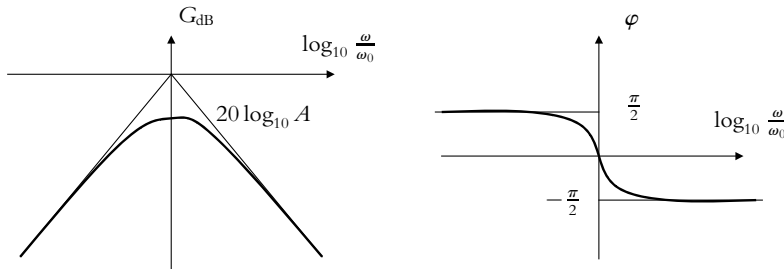
17. On applique la relation fournie soit $\frac{N_2}{N_1} = \frac{U_2}{U_1} = 2,17 \cdot 10^{-2}$.
18. On en déduit que $N_2 = 11$ spires.
19. La puissance consommée par le capteur est

$$P_c = \langle -u_2(t)i_2(t) \rangle = \langle \frac{u_2^2}{R_e} \rangle = \frac{N_2^2}{N_1^2 R_e} \langle u_1^2(t) \rangle = \frac{N_2^2 V_M^2}{2N_1^2 R_e} = 2,45 \text{ mW}$$

Cette valeur est négligeable devant la puissance mesurée donc il n'y a pas de perturbation.

B.4 1. On associe en série deux impédances respectivement $\underline{Z}_{||} = \frac{R}{1 + jRC\omega}$ et $\underline{Z}_{\perp} = \frac{1 + jRC\omega}{jC\omega}$. On obtient la fonction de transfert à partir de la relation du diviseur de tension soit $\underline{H} = \frac{\underline{\varepsilon}}{\underline{\varepsilon}} = \frac{\underline{Z}_{||}}{\underline{Z}_{\perp} + \underline{Z}_{||}} = \frac{jRC\omega}{1 + 3jRC\omega - R^2 C^2 \omega^2}$. On peut identifier terme à terme soit $A = Q = \frac{1}{3}$ et $\omega_0 = \frac{1}{RC}$.

2. Pour ω tendant vers 0, on a $\underline{H} \simeq j \frac{A\omega}{Q\omega_0}$ soit un gain en décibels $G_{dB} \simeq 20 \log_{10} \frac{\omega}{\omega_0}$.
Pour ω infini, on a $\underline{H} \simeq -j \frac{A\omega_0}{Q\omega}$ soit $G_{dB} \simeq -20 \log_{10} \frac{\omega}{\omega_0}$.
3. Le déphasage est tel que $\varphi = -\psi$ avec $\tan \psi = Q \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)$ et $\cos \psi > 0$ donc $\psi \in \left[-\frac{\pi}{2}, \frac{\pi}{2} \right]$ et $\varphi = -\text{Arctan} Q \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)$.
Pour ω tendant vers 0, le déphasage φ tend vers $\frac{\pi}{2}$.
Pour ω infini, le déphasage φ tend vers l'infini, on a l'expression $\varphi \rightarrow -\frac{\pi}{2}$.
4. On étudie les variations à partir de la fonction $f(\omega) = 1 + Q \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right) \geq 1$ pour tout ω et on a l'égalité pour $\omega = \omega_0$. On en déduit que le gain G est une fonction croissante puis décroissante. On a donc un phénomène de résonance pour $\omega = \omega_0$.
Pour le déphasage, on étudie la fonction $g(\omega) = \frac{\omega}{\omega_0} - \frac{\omega_0}{\omega}$. Sa dérivée vaut $g'(\omega) = \frac{1}{\omega_0} + \frac{\omega_0}{\omega^2} > 0$ donc g est croissante et $\varphi = -\text{Arctang}(\omega)$ décroissante.
On en déduit l'allure :



5. En utilisant le passage à l'équation différentielle par équivalence $j\omega$ et $\frac{d}{dt}$, on obtient l'équation différentielle $Q\frac{d^2s}{dt^2} + \omega_0\frac{ds}{dt} + Q\omega_0^2s = A\omega_0\frac{de}{dt}$.
6. Un amplificateur opérationnel est idéal si les courants de polarisation sont nuls (ou de façon équivalente si l'impédance d'entrée est infinie), si la résistance de sortie est nulle et si le gain différentiel statique est infini.
Un amplificateur opérationnel fonctionne en régime linéaire si on a la relation $s = \mu_0\epsilon$ avec $\epsilon = v_+ - v_-$.
7. En appliquant la relation du pont diviseur de tension, on a $v_- = \frac{R_1 e}{R_1 + R_2}$.
8. Comme l'amplificateur opérationnel est idéal et fonctionne en régime linéaire, on a $\epsilon = 0$ soit $v_- = s$.
Or $\frac{d^2s}{dt^2} + \frac{\omega_0}{Q} \left[1 - A \left(1 + \frac{R_2}{R_1} \right) \right] \frac{ds}{dt} + \omega_0^2 s = 0$
9. On a une sortie purement sinusoïdale si $1 - A \left(1 + \frac{R_2}{R_1} \right) = 0$ soit $A = \frac{R_1}{R_1 + R_2}$.
10. Si $1 - A \left(1 + \frac{R_2}{R_1} \right) < 0$, tous les coefficients ne sont pas de même signe : le système est instable et s tend vers l'infini. On en déduit le fonctionnement en saturation.

B.5 1. On étudie le système constitué de la charge q . On se place dans le référentiel du laboratoire galiléen. Le bilan des forces est le poids qui sera négligé, la force de frottement fluide $-\alpha \vec{v}$, la force électrique $q\vec{E}$ et par application du principe fondamental de la dynamique $m\frac{d\vec{v}}{dt} = -\alpha \vec{v} + q\vec{E}$ soit $\frac{d\vec{v}}{dt} + \frac{\alpha}{m}\vec{v} = \frac{q}{m}\vec{E}$.

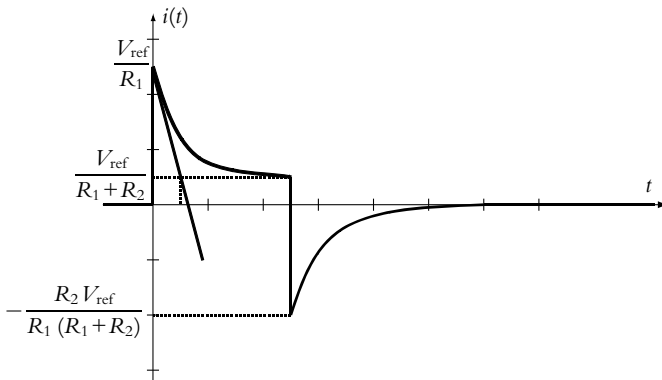
2. La solution est la somme de la solution générale décrivant régime transitoire $\vec{V} \exp\left(-\frac{\alpha}{m}t\right)$ de temps caractéristique $\tau = \frac{m}{\alpha}$ et d'une solution particulière décrivant régime permanent $\frac{q}{\alpha}\vec{E}$ qui est la vitesse limite $\vec{v}_l$ soit

$$\vec{v} = \vec{V} \exp\left(-\frac{\alpha}{m}t\right) + \frac{q}{\alpha}\vec{E}$$

3. En appliquant la définition, on obtient $\vec{j} = \frac{nq^2}{\alpha}\vec{E}$.
4. En identifiant ce résultat à la loi d'Ohm et en utilisant la relation $\alpha = \frac{m}{\tau}$, on exprime la conductivité $\sigma = \frac{nq^2\tau}{m}$.
5. On calcule le travail élémentaire $\delta W = qE\vec{u}_x \cdot d\vec{OM} = d(qEx) = -d(qV)$. On a donc une force conservative dont le potentiel s'écrit $V = -Ex$ par identification.
6. La tension s'écrit $U = V(L) - V(0) = -EL$ et comme le courant est lié aux électrons $I = -\iint \vec{j} \cdot d\vec{S} = -\sigma ES$. On en déduit la valeur de la résistance $R = \frac{U}{I} = \frac{1}{\sigma} \frac{L}{S}$.
7. Un amplificateur opérationnel est idéal si les courants de polarisation sont nuls (ou de façon équivalente si l'impédance d'entrée est infinie), si la résistance de sortie est nulle et si le gain différentiel statique est infini.
8. Un amplificateur opérationnel fonctionne en régime linéaire si on a la relation $s = \mu_0\epsilon$ avec $\epsilon = v_+ - v_-$.
9. La sortie de l'amplificateur opérationnel ne prend pas de valeur infinie, ce qui implique avec l'hypothèse que μ_0 est infini que ϵ doit être nul.

10. On a $v_+ = V_{\text{ref}}$ et en utilisant le théorème de Millman $v_- = \frac{I + \frac{v_0}{R}}{\frac{1}{R}} = RI + v_0$. Compte tenu des hypothèses, $\epsilon = 0$ et $v_0 = V_{\text{ref}} - RI$.
11. La relation du pont diviseur de tension donne $v_+ = \frac{R'_2 e_2}{R'_1 + R'_2}$ et l'application du théorème de Millman donne $v_- = \frac{\frac{e_1}{R_1} + \frac{s}{R_2}}{\frac{1}{R_1} + \frac{1}{R_2}} = \frac{R_2 e_1 + R_1 s}{R_1 + R_2}$ soit avec l'hypothèse $\epsilon = 0$, on en déduit $s = \frac{R'_2(R_1 + R_2)}{R_1(R'_1 + R'_2)} e_2 - \frac{R_2}{R_1} e_1$.
12. Si $R_1 = R_2 = R'_1 = R'_2 = R$, on peut écrire $s = e_2 - e_1$. On a donc un soustracteur.
13. On a $v_+ = v_1$ et en appliquant la relation du pont diviseur de tension, on obtient $v_- = \frac{R\nu}{R + R_p}$. Compte tenu des hypothèses, on a $\epsilon = 0$ et $\nu = \left(1 + \frac{R_p}{R}\right) v_1$. Il s'agit d'un amplificateur non inverseur.
14. La loi des mailles donne $u + R_1 i = e$ et la loi des nœuds $i = \frac{u}{R_2} + i_C = \frac{u}{R_2} + C \frac{du}{dt}$. Finalement $\left(1 + \frac{R_1}{R_2}\right) u + R_1 C \frac{du}{dt} = e$.
15. La tension u est continue du fait de la continuité de la tension aux bornes d'une capacité donc $u(0^+) = u(0^-) = 0$. En reportant dans l'équation différentielle précédente, on en déduit $\frac{du}{dt}(0^+) = \frac{V_{\text{ref}}}{R_1 C}$. On reporte ce résultat dans la loi des nœuds et on obtient $i(0^+) = \frac{u(0^+)}{R_2} + C \frac{du}{dt}(0^+) = \frac{V_{\text{ref}}}{R_1}$.
16. La solution est la somme de la solution générale $U \exp\left(-\frac{(R_1 + R_2)t}{R_1 R_2 C}\right)$ et d'une solution particulière constante $\frac{R_2 V_{\text{ref}}}{R_1 + R_2}$. On détermine la constante U à partir de la condition initiale $u(0) = 0 = U + \frac{R_2 V_{\text{ref}}}{R_1 + R_2}$. Finalement $u = \frac{R_2 V_{\text{ref}}}{R_1 + R_2} \left(1 - \exp\left(-\frac{(R_1 + R_2)t}{R_1 R_2 C}\right)\right)$.
17. Si le régime permanent est atteint à Δt , l'exponentielle est asymptotiquement nulle soit $u(\Delta t^-) = \frac{R_2 V_{\text{ref}}}{R_1 + R_2}$. D'autre part, $i_C(\Delta t^-) = 0$ donc $i(\Delta t^-) = \frac{u(\Delta t^-)}{R_2} + i_c(\Delta t^-) = \frac{V_{\text{ref}}}{R_1 + R_2}$.
18. On obtient i à partir de $i = \frac{u}{R_2} + C \frac{du}{dt} = \frac{V_{\text{ref}}}{R_1 + R_2} \left(1 + \frac{R_2}{R_1} \exp\left(-\frac{(R_1 + R_2)t}{R_1 R_2 C}\right)\right)$.
19. Il faut résoudre l'équation différentielle $\left(1 + \frac{R_1}{R_2}\right) u + R_1 C \frac{du}{dt} = 0$ soit $u = U' \exp\left(-\frac{(R_1 + R_2)t}{R_1 R_2 C}\right)$ et par continuité en $t = \Delta t$ de la tension aux bornes de la capacité $U' \exp\left(-\frac{(R_1 + R_2)t}{R_1 R_2 C}\right) = \frac{R_2 V_{\text{ref}}}{R_1 + R_2}$. Finalement la solution s'écrit $u = \frac{R_2 V_{\text{ref}}}{R_1 + R_2} \exp\left(-\frac{(R_1 + R_2)(t - \Delta t)}{R_1 R_2 C}\right)$.
20. On obtient i à partir de $i = \frac{u}{R_2} + C \frac{du}{dt} = -\frac{R_2 V_{\text{ref}}}{R_1 (R_1 + R_2)} \exp\left(-\frac{(R_1 + R_2)t}{R_1 R_2 C}\right)$.

21. En utilisant la question 10, on a $v_0 = V_{\text{ref}} - R_f i(t)$.
22. En utilisant la question 12, on en déduit $v_1 = V_{\text{ref}} - v_0 = R_f i(t)$.
23. En utilisant la question 13, on obtient $v = \left(1 + \frac{R_v}{R}\right) R_f i(t)$.
24. En utilisant la réponse à la question 6 avec $L = h$ et $S = \pi \left(\frac{d}{2}\right)^2$ soit une résistance $R_d = \frac{4h}{\sigma \pi d^2} = 2,5 \text{ M}\Omega$.
25. On fait la même chose avec $L = h$ et $S = 2\pi \frac{d}{2} e$ soit $R_s = \frac{h}{\sigma \pi d e} = 2,1 \text{ G}\Omega$.
26. La résistance R_p est en série avec l'association en parallèle de R_d et R_s . On en déduit après $\frac{R_d R_s}{R_d + R_s}$.
On a donc une résistance équivalente $R_{\text{eq}} = R_p + \frac{R_d R_s}{R_d + R_s} \simeq R_p + R_d$.
27. D'après question 23, la tension $v = \left(1 + \frac{R_v}{R}\right) R_f i(t) = R_f i(t)$ avec $R_v = 0 \text{ }\Omega$ et d'après les questions 18 et 20 allure de $i(t)$ suivante :



On en déduit $\frac{R_f}{R_1} V_{\text{ref}} = 0,24 \text{ V}$ soit $R_1 = 2,08 \text{ M}\Omega$.

On a $\frac{R_f}{R_1 + R_2} V_{\text{ref}} = 0,04 \text{ V}$ soit $R_1 + R_2 = 12,5 \text{ M}\Omega$ et $R_2 = 10,4 \text{ M}\Omega$.

La valeur de τ est donnée par l'intersection de la tangente à l'origine et de l'asymptote à l'infini soit $0,1 \text{ s}$ donc $\frac{R_1 R_2 C}{R_1 + R_2} = 0,1 \text{ s}$ soit $C = 0,58 \text{ nF}$.

Chapitre 22

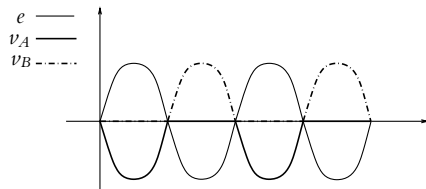
- A.1** 1. Il faut résoudre le système formé par $U = E - RI$ et $I = \frac{1}{R_d}(U - U_0)$. On obtient $U = \frac{R_d E + R U_0}{R_d + R} = 1,67 \text{ V}$ et $I = \frac{E - U_0}{R_d + R} = 21,6 \text{ mA}$.
2. Une variation ΔE provoque une variation $\Delta U = \frac{R_d \Delta E}{R_d + R}$. Le facteur de régulation s'écrit donc $f_0 = \frac{R_d}{R_d + R} = 6$ et le taux d'ondulation $\frac{\Delta U}{U} = \frac{R_d \Delta E}{R_d E + R U_0} = 0,1$.

3. L'ensemble constitué de la résistance R_C et de l'alimentation admet comme modèle de Thévenin équivalent un générateur de force électromotrice $E' = \frac{R_C E}{R_C + R}$ et de résistance $R' = \frac{R R_C}{R + R_C}$. La diode est passante si $I > 0$. Il faut donc que $E' > U_0$ et $R_C > \frac{R U_0}{E - U_0}$. Le courant vérifie alors $\frac{E' - U_0}{R_d + R} < I_{\max}$ qu'on peut écrire $R_C < \frac{R(U_0 + (R + R_d) I_{\max})}{E - U_0 - (R + R_d) I_{\max}}$. L'application numérique donne $30,8 \, \Omega < R_C < 800 \, \Omega$.

- A.2** 1. Si la diode est passante, on la remplace par l'association en série d'une source idéale de tension de force électromotrice U_S et une résistance R_d . On applique alors le théorème de Millman usuel sur ce nouveau circuit et on en déduit $i = \frac{R_2 R_d E_1 - R_1 R_d E_2 + R_1 R_2 (E_3 - U_S)}{R R_1 R_2 + R R_1 R_d + R R_2 R_d + R_1 R_2 R_d}$.
2. Lorsque la diode est bloquée, on supprime la branche contenant la diode puis on applique le théorème de Millman usuel. On obtient ici $i = \frac{R_2 E_1 - R_1 E_2}{R_1 R_2 + R R_1 + R R_2}$.

On peut noter qu'ici le théorème de Millman réservé aux circuits linéaires peut être appliqué grâce au fait que la non linéarité peut se décomposer en deux cas linéaires : on a un circuit "linéaire par morceaux".

- A.3** 1. Si la diode est bloquée, on la remplace par un interrupteur ouvert donc $i = 0$ et $u < U_S$. Si la diode est passante, elle est modélisée par l'association en série d'une résistance R_d et d'une source idéale de tension de force électromotrice U_S . On a $u = U_S + R_d i$ et $u \geq U_S$.
2. La diode A est bloquée si $v_A - s < U_S$ et la diode B est bloquée si $s - v_B < U_S$.
3. a) L'amplificateur opérationnel étant idéal et en régime linéaire, on a ici $v_+ = v_- = 0$. L'application du théorème de Millman en v_- donne $s = U_S - \frac{R + R_d}{R} e$. On en déduit $v_A = 0$ et $v_B = -e$.
- b) On doit vérifier les conditions correspondant à ce cas soit $e < \frac{2R U_S}{R + R_d}$ et $e < 0$, ce qui se résume à $e < 0$.
4. a) L'amplificateur opérationnel étant idéal et en régime linéaire, on a ici $v_+ = v_- = 0$. L'application du théorème de Millman en v_- donne $s = -U_S - \frac{R + R_d}{R} e$. On en déduit $v_B = 0$ et $v_A = -e$.
- b) On doit vérifier les conditions correspondant à ce cas soit $e > 0$ et $e > -\frac{2R U_S}{R + R_d}$, ce qui se résume à $e > 0$.
5. Il n'y a pas d'autres cas car toutes les valeurs de e ont été envisagées.
6. On en déduit l'allure suivante pour les signaux :



- A.4** 1. Deux lois des mailles donnent :

$$\begin{cases} e = -u_1 - u \\ e = u_2 + u \end{cases}$$

2. On remarque que :

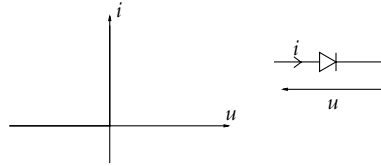
- pour $0 < t < \frac{T}{2}$, $u_1 < 0$ et $u_2 = 0$ donc D_2 est passante et D_1 est bloquée ;
- pour $\frac{T}{2} < t < T$, $u_2 < 0$ et $u_1 = 0$ donc D_1 est passante et D_2 est bloquée.

3. Ainsi d'après les relations précédentes, pour $0 < t < \frac{T}{2}$, $u_2 = 0$ donc $u = e > 0$. Pour la diode D_1 , on a donc $u_1 = -e - u = -2e < 0$ et D_1 est bien bloquée.

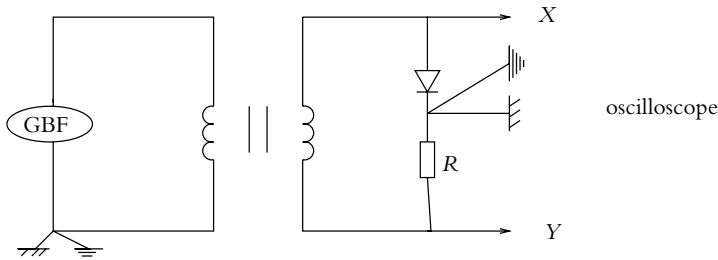
Lorsque à $t = \frac{T}{2}$, e devient négatif, D_2 se bloque et D_1 devient passante alors $u = -e > 0$. Dans tous les cas, i et u sont positifs : on a réalisé un redressement double-alternance.

Ce circuit est plus simple qu'un pont de Graetz mais l'inconvénient est que les diodes doivent supporter une tension inverse égale à $2e$ en valeur absolue.

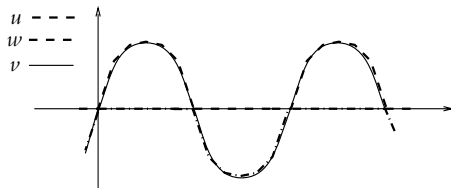
- B.1** 1. a) La caractéristique d'une diode idéale est :



- b) Pour obtenir la caractéristique à l'oscilloscope, il faut utiliser celui-ci en mode XY, enclencher $-CH2$ pour avoir i et non $-i$ et utiliser un GBF en sinusoïdal dans le montage suivant :



- c) Le transformateur d'isolement permet de s'affranchir de la masse du GBF reliée à la Terre et de pouvoir placer la masse de l'oscilloscope où on veut.
- d) L'écriture de la loi des mailles donne $v + Ri + u = 0$ donc la diode est bloquée si $u \leq 0$. Dans ce cas, on a $i = 0$, $u = v$ et $w = -Ri = 0$ pour $v \leq 0$. Sinon la diode est passante et $u = 0$ et $w = v = -Ri$. On obtient les allures suivantes pour les différents signaux :



- e) Si la tension est constante au primaire, on a $v = 0$ et $i = 0$.

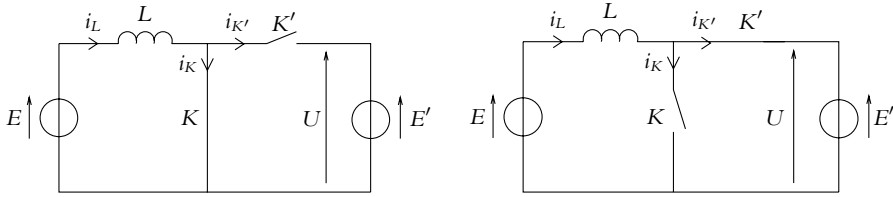
2. Cas $u = E'$:

- a) Pour $0 \leq t \leq \alpha T$, on a $i_{K'} = 0$, $i_L = i_K$ et $E = L \frac{di_L}{dt}$. Par intégration, on obtient

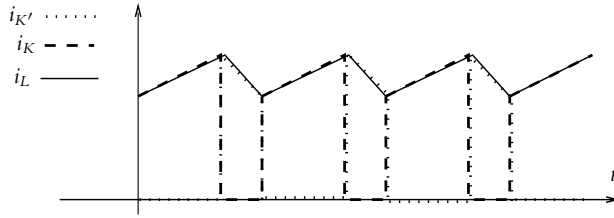
$$i_L = i_K = I_m + \frac{E}{L}t \text{ et } i_{K'} = 0.$$

Pour $\alpha T \leq t \leq T$, on a $i_K = 0$, $i_L = i_{K'}$ et $E = E' + L \frac{di_L}{dt}$. Par intégration, on obtient $i_K = 0$

$$\text{et } i_L = i_{K'} = \frac{E - E'}{L}t + I_m + \frac{E'\alpha T}{L}.$$



b) On en déduit les allures suivantes pour les signaux :



c) L'intensité i_L est périodique donc $i_L(0) = i_L(T)$, ce qui se traduit par $E = E'(1 - \alpha)$.

d) $I_M = i_L(T)$ donc $\Delta i_L = \frac{E\alpha T}{L} \leq \Delta i_{L,max}$ dont on déduit $L \geq \frac{E\alpha T}{\Delta i_{L,max}} = 5,0 \text{ mH}$.

e) La puissance s'obtient par $P_g = \frac{1}{T} \int_0^T E i_L(t) dt = E \left(\frac{\alpha ET}{2L} + I_m \right)$.

f) On en déduit $I_m = \frac{P_g}{E} - \frac{\alpha ET}{2L} = 2,85 \text{ A}$ et $I_M = I_m + \Delta i_{L,max} = 3,15 \text{ A}$.

g) Les portions de caractéristiques décrites sont les suivantes :



h) L'interrupteur K est un transistor et K' une diode.

i) La valeur moyenne de la tension aux bornes de l'interrupteur K est

$$V_0 = \frac{1}{T} \int_0^T u_K(t) dt = E$$

3. a) La valeur moyenne du courant dans la capacité C est

$$I_c = \frac{1}{T} \int_0^T dq = 0$$

car $q(t)$ est périodique de période T .

b) De même, la valeur moyenne du courant dans la résistance R est

$$I_R = \frac{1}{T} \int_0^T i_R(t) dt = \frac{1}{T} \int_0^T (i_{K'}(t) - i_C(t)) dt = \frac{P_g(1 - \alpha)}{E}$$

c) La puissance moyenne dissipée dans la capacité est

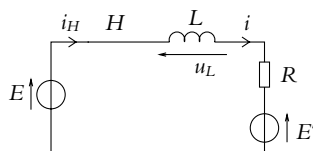
$$P_C = \frac{1}{T} \int_0^T u_C(t) i_C(t) dt = \frac{C}{T} \int_0^T u_C du_C = 0$$

car la tension u_C est périodique de période T .

d) On établit de même que $P_R = P_g$.

e) Tout se passe donc comme si u était constante. Par conséquent, l'étude précédente avec $u = E'$ est tout à fait justifiée.

- B.2** 1. a) Lorsque l'interrupteur est fermé, la tension u_D aux bornes de la diode est égale à $-E < 0$: la diode est bloquée. Le schéma équivalent au circuit est celui de la figure ci-contre.



b) On a $i = i_H$ et l'équation différentielle s'écrit :

$$\frac{di_H}{dt} + \frac{R}{L}i_H = \frac{E - E'}{L}$$

La solution générale de cette équation est :

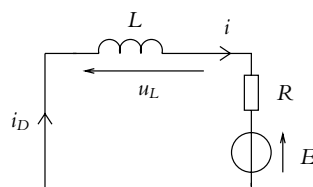
$$i_H(t) = i(t) = A \exp\left(-\frac{t}{\tau}\right) + \frac{E - E'}{R}$$

où A désigne une constante et $\tau = \frac{L}{R}$ la constante de temps caractéristique du circuit. Cette solution est valable pour $0 < t < \alpha T$

► **Remarque** : vues les valeurs numériques de E et de E' , c'est le générateur E qui impose le sens du courant.

2. Juste après l'ouverture de l'interrupteur, le courant étant continu dans la bobine, $i(t)$ est positif et la diode devient passante. La diode permet au courant $i(t)$ de ne pas s'annuler lors de l'ouverture de H .

Le schéma équivalent du circuit est donné sur la figure ci-contre.



a) Dans ce cas, $i_D = i$ et l'équation différentielle du circuit s'écrit :

$$\frac{di_D}{dt} + \frac{R}{L}i_D = -\frac{E'}{L}$$

La solution générale de cette équation est :

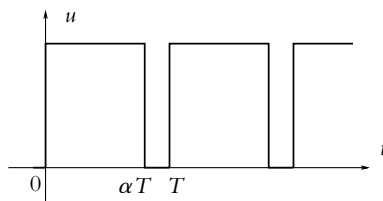
$$i_D(t) = i(t) = A' \exp\left(-\frac{t}{\tau}\right) - \frac{E'}{R}$$

où A' désigne une constante et $\tau = \frac{L}{R}$ la constante de temps du circuit (on note que c'est la même que dans l'état précédent). Cette solution est valable pour $\alpha T < t < T$.

3. a) Lorsque l'interrupteur est fermé, $u(t) = E$ et lorsqu'il est ouvert, $u(t) = 0$, ce qui donne l'évolution ci-contre. La valeur moyenne est :

$$\langle u \rangle = \frac{1}{T} \int_0^T u(t) dt = \frac{1}{T} \int_0^{\alpha T} E dt = \alpha E$$

La composante continue est donc αE .

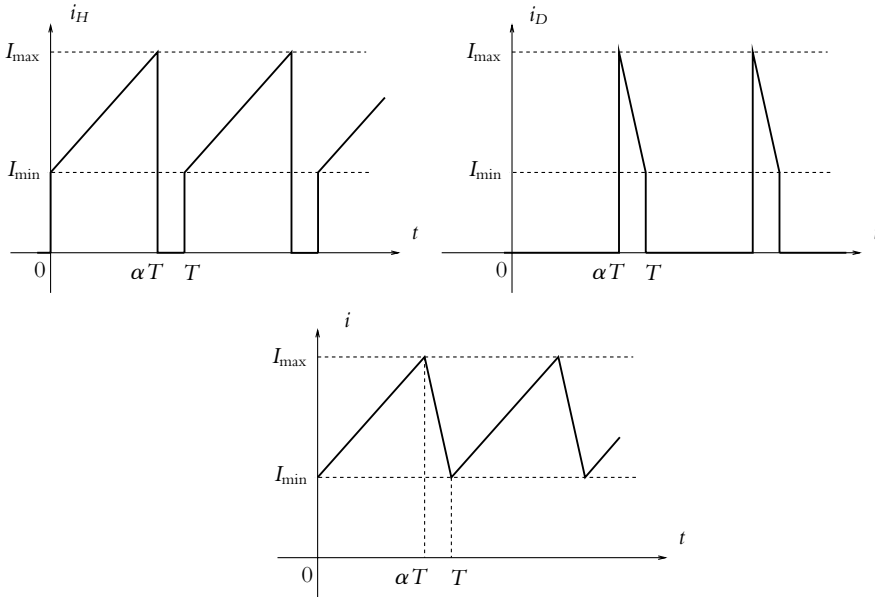


- b) Le calcul numérique donne $\tau = \frac{L}{R} = 4 \cdot 10^{-2}$ s soit $\tau = 40T$. On peut donc considérer dans les expressions des intensités que $t \ll \tau$ et effectuer un développement limité au premier ordre, ce

qui donne :

$$\begin{cases} \text{pour } 0 < t < \alpha T : i_H = A \left(1 - \frac{t}{\tau}\right) + \frac{E - E'}{R} \text{ et } i_D = 0 \\ \text{pour } \alpha T < t < T : i_D = A' \left(1 - \frac{t}{\tau}\right) - \frac{E'}{R} \text{ et } i_H = 0 \end{cases}$$

c) L'allure des courants i_H et i_D est obtenue grâce aux expressions précédentes et $i(t) = i_H(t) + i_D(t)$.



On remarquera que seul le courant $i(t)$ est une fonction continue du temps : cette propriété est imposée par la présence de la bobine.

d) La tension aux bornes de la bobine est :

$$u_L = L \frac{di}{dt}$$

Sa valeur moyenne est donc :

$$\langle u_L \rangle = \frac{1}{T} \int_0^{\alpha T} L \frac{di}{dt} dt + \frac{1}{T} \int_{\alpha T}^T L \frac{di}{dt} dt = \frac{L}{T} (I_{\max} - I_{\min} + I_{\min} - I_{\max}) = 0$$

Pour calculer la valeur moyenne de i , on écrit que :

$$\begin{cases} \text{pour } 0 < t < \alpha T : E = u_L + E' + Ri \\ \text{pour } \alpha T < t < T : 0 = u_L + E' + Ri \end{cases}$$

soit :

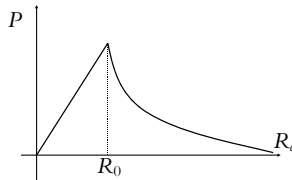
$$i(t) = \frac{1}{T} \int_0^T i dt = \frac{1}{RT} \left(\int_0^{\alpha T} E dt - \int_0^T E' dt \right) + \frac{\langle u_L \rangle}{R} = \frac{\alpha E - E'}{R}$$

e) Lors de l'ouverture de l'interrupteur, la bobine crée une force électromotrice pour assurer la continuité du courant : c'est ce qui rend la diode passante. Sans la bobine, le courant serait nul entre αT et T .

B.3 1. On a $i'(t) = \frac{I_0}{V_0} e^{\frac{u}{V_0}} > 0$ donc la fonction i est croissante de $-I_0 - I_p$ à $+\infty$.

2. Il faut résoudre $i(u) = 0$. On obtient $u_{C0} = V_0 \ln \left(1 + \frac{I_p}{I_0} \right) = 0,102 V$.

3. La puissance reçue est $P = ui$ donc la puissance fournie est positive si $ui < 0$ soit ici $u > 0$ et $i < 0$.
4. On a une source idéale de courant si le courant traversant la source est constant pour toute tension u appliquée. C'est le cas ici si $u < 0$.
5. On a la relation $u = -R_c i$.
6. Le premier cas est $-R_c i = u_{C0}$ soit $u = u_{C0}$ et $i = -\frac{u_{C0}}{R_c}$ pour $i > -I_0 - I_p$. On a donc $R_c > \frac{u_{C0}}{I_0 + I_p}$.
De même, pour le second cas, on a pour $R_c < \frac{u_{C0}}{I_0 + I_p}$ une intensité $i = -I_0 - I_p$ et une tension $u = R_c (I_0 + I_p)$.
7. La puissance fournie est $P = -ui$ donc $P = \frac{u_{C0}^2}{R_c}$ pour $R_c > R_0$ avec $R_0 = \frac{u_{C0}}{I_0 + I_p}$ et $P = R_c (I_0 + I_p)^2$ pour $R_c < R_0$.
8. La puissance en fonction de R_c a l'allure suivante :



9. On a un maximum de la puissance pour $R_{opt} = R_c = R_0$ soit $P_{max} = u_{C0} (I_0 + I_p)$.
10. Le rendement s'écrit $\eta = \frac{P_{max}}{P_l} = kV_0 \left(1 + \frac{1}{x}\right) \ln(1+x)$.
11. L'application numérique donne $\eta = 0,052$.
12. Pour x infini, le rapport $\frac{I_p}{I_0}$ est infini comme u_{C0} et on a un générateur de courant idéal $i = -I_p$. Il y a donc un problème avec le modèle.
13. La puissance maximale P_{max} est telle que l'aire du rectangle entre $i = 0$ et $i = -I$ et $u = 0$ et $u = U$ avec I et U valeurs permettant l'obtention de P_{max} donc on a une surestimation lorsqu'on utilise des droites avec I et U plus grandes.
14. Du fait de l'association en séries des résistances, on peut écrire

$$u'_{C0} = Nu_{C0} = NV_0 \ln \left(1 + \frac{kP_l}{I_0}\right)$$

$$\text{et } i'_{CC} = -I_p - I_0 = -(kP_0 + I_0).$$

15. On en déduit $P'_{max} = NP_{max}$.
16. De même, on a $R'_{opt} = NR_{opt}$.
17. Pour une association en parallèle, on a $u''_{C0} = u_{C0}$, $i''_{CC} = -N(I_0 + I_p)$, $P''_{max} = NP_{max}$ et $R''_{opt} = \frac{R_{opt}}{N}$.
18. On a $R_c = 5R_{opt}$ donc il faut 5 photodiodes en série et $\eta' = 0,26$.
19. En statique, la capacité est équivalente à un interrupteur ouvert donc on remplace photodiode par l'association en parallèle de sources idéales de courant I_p et de R_d soit par loi des mailles après passage du modèle de Norton au modèle de Thévenin $v_s = \frac{R}{R + R_d} (E - kR_d P_l)$.
20. La tension v_s est proportionnelle à P_l si $E \ll kR_d P_l$ et $A = k \frac{RR_d}{R + R_d} = 4995$.
21. On a $u = E - v_s$ avec $u < 0$ donc $P_l < P_{sat} = -\frac{E}{A} = 3,0 \text{ mW}$.

22. On en déduit $\sigma = \frac{|v_{SN}|}{A} = 2, 0 \cdot 10^{-4} \text{ mW}$.

23. En régime non continu, on garde C_d . En appliquant la loi des mailles et la loi des nœuds, on obtient $i + I_p = i_1 + i_2$ en notant i_1 et i_2 les courants respectifs dans R_d et C_d soit $i_2 = R_d C_d \frac{di_1}{dt}$ et $E = R_d i_1 + R_i$ donc après calculs, on obtient $\frac{dv_s}{dt} + \frac{v_s}{\tau} = \frac{E}{R_d C_d} - \frac{A}{\tau} P_l$ avec $\tau = \frac{R R_d C_d}{R + R_d}$.

24. Pour $E \ll k R_d P_l$, l'équation différentielle devient

$$\frac{dv_s}{dt} + \frac{v_s}{\tau} \simeq -\frac{A}{\tau} (P_0 + \Delta P \cos \omega t)$$

On en déduit $v_s = V \exp\left(\frac{t}{\tau}\right) - A P_0 + v_{\text{sinus}}(t)$ où v_{sinus} est la solution de l'équation différentielle $\frac{dv_s}{dt} + \frac{v_s}{\tau} = -\frac{A}{\tau} \Delta P \cos \omega t$ donc $v_{s0} = -A P_0$.

25. En prenant la notation complexe, on a $j\omega \underline{\Delta v_s} + \frac{1}{\tau} \underline{\Delta v_s} = -\frac{A}{\tau} \underline{\Delta P_l}$ et $\underline{H} = \frac{-A}{1 + j\omega \tau}$.

26. Le gain G est une fonction décroissante de $G_{\text{dB}} = 20 \log_{10} A$ pour ω nul et $G_{\text{dB}} = 20 \log_{10} A - 20 \log_{10} \omega \tau$ avec une pente à -20 dB par décade pour ω infini.

Le déphasage s'exprime par $\varphi = \pi - \text{Arctan} \omega \tau$ qui est une fonction décroissante de π à $\frac{\pi}{2}$.

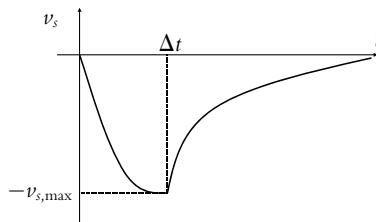
27. La pulsation de coupure s'écrit pour ce montage $\omega_c = \frac{1}{\tau} = 50, 1 \cdot 10^6 \text{ rad.s}^{-1}$.

28. Si $0 \leq t \leq \Delta t$, il faut résoudre $\frac{dv_s}{dt} + \frac{v_s}{\tau} = -\frac{A}{\tau} P_0$ soit

$$v_s = -A P_0 \left(1 - \exp\left(-\frac{t}{\tau}\right)\right)$$

Pour $t \geq \Delta t$, il faut résoudre $\frac{dv_s}{dt} + \frac{v_s}{\tau} = 0$. On obtient

$$v_s = -A P_0 \left(1 - \exp\left(-\frac{\Delta t}{\tau}\right)\right) \exp\left(-\frac{t - \Delta t}{\tau}\right)$$



29. On doit résoudre $v_s(t) = -\frac{v_{s,\text{max}}}{2}$ soit $t' = -\tau \ln \frac{1 + \exp\left(-\frac{\Delta t}{\tau}\right)}{2}$ et $t'' = \Delta t + \tau \ln 2$. On en déduit

$$\Delta t_d = \Delta t + \tau \ln 2 + \tau \ln \frac{1 + \exp\left(-\frac{\Delta t}{\tau}\right)}{2}.$$

30. On a $\Delta t_d \simeq \Delta t$ si $\Delta t \gg \tau$. De même, $\Delta t_d \simeq \tau \ln 2$ si $\Delta t \ll \tau$.

31. On a un extremum pour $\frac{d\Delta t_d}{d\Delta t} = \frac{1 - \exp\left(-\frac{\Delta t}{\tau}\right)}{1 + \exp\left(-\frac{\Delta t}{\tau}\right)}$. On l'obtient pour $\Delta t = 0$ et la quantité est

négative donc la plus petite valeur est obtenue pour $\Delta t = 0$ soit $\tau \ln 2$.

32. Il n'y a pas de recouvrement si $\frac{1}{B} > \tau \ln 2$ donc il faut $B \simeq \frac{1}{\tau \ln 2}$.
33. On a $\frac{B}{\sigma} = \frac{k}{C_d \ln 2 |v_{SN}|}$ avec k et C_d caractéristiques du détecteur.
34. On veut $B > 10^7 \text{ s}^{-1}$ soit $\sigma > \frac{BC_d \ln 2 |v_{SN}|}{k} = 2,78 \cdot 10^{-8} \text{ W}$.

B.4 1. a) Avec les données numériques, on calcule $I_{ph} = 0,14 \text{ A}$.

La caractéristique est représentée avec une convention récepteur, donc la puissance reçue par la diode est $P = ui$. Pour que la diode fonctionne en générateur il faut que P soit négative, ce qui correspond aux quadrants II et IV. Comme il n'y a rien dans le quadrant II, le fonctionnement photovoltaïque a lieu dans le quadrant IV.

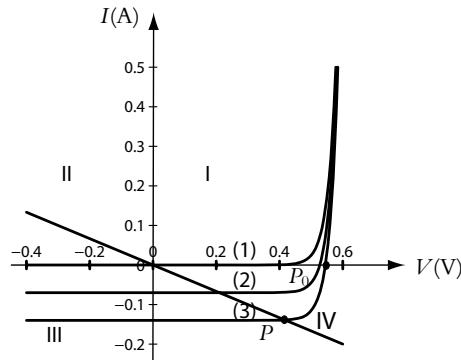


Figure S22.1

Si la photodiode fonctionne en générateur, elle alimente une résistance R par exemple (figure S22.1). On peut alors écrire $V = -RI$ et le point de fonctionnement P est l'intersection de la caractéristique (courbe (3) d'après la valeur calculée précédemment) avec la droite $V = -RI$ (figure S22.1).

- b) En circuit ouvert, cela revient à avoir $i = 0$ où $R \rightarrow \infty$, donc le point de fonctionnement est P_0 à l'intersection de la caractéristique avec l'axe des abscisses. On trouve une tension de l'ordre de $0,5 \text{ V}$.

Cette valeur limite est obtenue pour $I = 0$ soit $I_d = I_s [\exp(\frac{qV}{kT}) - 1] - I_{ph} = 0$. Avec $I_{ph} \gg I_s$, on doit résoudre $I_{ph} = I_s \exp(\frac{qV}{kT})$ soit $V_{CO} = \frac{kT}{q} \ln \frac{I_{ph}}{I_s}$.

La puissance électrique est nulle puisque $I = 0$.

- c) En court-circuit, la tension v est nulle donc l'intensité est $I_{cc} = 0 - I_{ph}$. Comme précédemment, la puissance est nulle puisque cette fois-ci $v = 0$.
- d) Pour une puissance donnée P_M , la courbe $I = f(V)$ est une hyperbole $I = P_M / V$. Cette courbe doit avoir une intersection avec la caractéristique, et la plus grande puissance est obtenue lorsque l'hyperbole est tangente à la caractéristique (figure S22.2). L'intensité débitée est prise positive comme le précise l'énoncé et l'on trouve $P_M = V_M I_M$ avec environ $I_M = 0,13 \text{ A}$, $V_M = 0,5 \text{ V}$ et $P_M = 65 \text{ mW}$.

La résistance correspondante est $R_{opt} = V_M / I_M = 3,85 \Omega$.

- e) Le rendement est donc $\eta = P_M / (\Sigma \mathcal{E}) = 0,16$. Ce rendement est faible.

2. a) Pour une association en parallèle, on ajoute les intensités. Si l'on considère qu'une cellule débite I_M alors n_p cellules débitent $I_{MC} = n_p I_M = 0,78 \text{ A}$.
Pour une association en série, on ajoute les tensions. Si l'on considère qu'une cellule a une tension V_M alors n_s cellules auront une tension $V_{MC} = n_s V_M = 18 \text{ V}$.

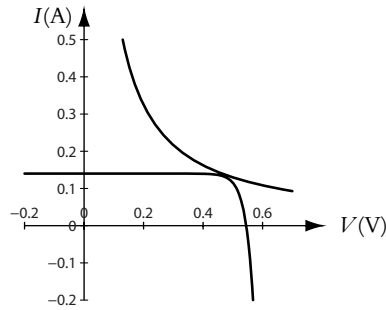


Figure S22.2

Pour déterminer la nouvelle caractéristique, il suffit d'ajouter les courants de court-circuit lorsque les cellules sont en parallèle et les tensions à vide lorsqu'elles sont en série. L'allure générale est la même que pour une seule cellule (figure S22.3).

- b) Comme on l'a signalé au début de l'exercice, si on branche une résistance R sur l'ensemble, on trouve le point de fonctionnement en traçant la caractéristique de la résistance qui est en convention générateur si la diode est en convention récepteur. La caractéristique de la résistance est $V = -RI$ (figure S22.3).

Pour la résistance $R_1 = 22 \, \Omega$, le point d'intersection F_1 donne $i_1 = 0,81 \, \text{A}$; $v_1 = 18 \, \text{V}$ et $P_1 = 14,6 \, \text{W}$.

Pour la résistance $R_2 = 10 \, \Omega$, le point d'intersection F_2 donne $i_2 = 0,82 \, \text{A}$; $v_2 = 8,2 \, \text{V}$ et $P_2 = 6,7 \, \text{W}$.

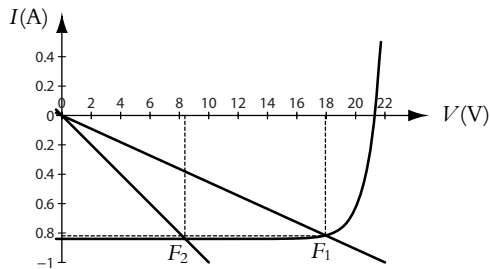


Figure S22.3

- c) Cela dépend s'il s'agit de charger une batterie ou d'alimenter une batterie en charge. Dans le deuxième cas, la batterie impose une tension (f.c.e.m), il n'y a plus qu'à tracer une verticale sur la caractéristique pour déterminer i .

B.5 1. On a représenté l'allure de la caractéristique de la diode sur la figure (S22.4).

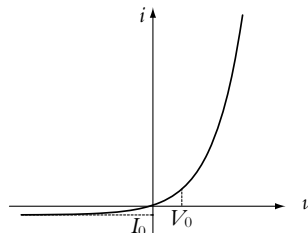


Figure S22.4

2. a) On note v^- la tension de l'entrée non inverseuse. On prend comme sens conventionnel du courant celui de la caractéristique de la diode. Puisqu'il n'y pas de courant pénétrant dans l'entrée non inverseuse, l'intensité i est celle traversant la résistance R soit $v_e - v^- = Ri$. En fonctionnement linéaire, on a $v^+ = v^-$ or ici $v^+ = 0$ d'où, $i = v_e/R$. D'autre part, la tension v aux bornes de la diode est égale à $v^- - v_s = -v_s$. Si l'on reporte v et i dans l'équation de la caractéristique, on trouve :

$$v_s = -v_0 \ln \left(1 + \frac{v_e}{RI_0} \right)$$

- b) Pour $0 < \omega t < \pi$, la tension v_e est positive, donc le logarithme est toujours défini et positif car l'argument est supérieur à 1. La tension v_s est négative. Peut-elle atteindre $-V_{\text{sat}} = -20$ V ?

$$v_0 \ln \left(1 + \frac{v_e}{RI_0} \right) > 20 \text{ V} \Rightarrow \ln \left(1 + \frac{v_e}{RI_0} \right) > 50 \Rightarrow \frac{v_e}{RI_0} > e^{50}$$

Pour qu'il y ait saturation il faut une tension v_e trop importante.

Pour $\pi < \omega t < 2\pi$, la tension v_e est négative. Comme $V_e\sqrt{2} > RI_0$, l'argument du logarithme peut devenir nul et s'il n'y avait pas saturation alors $v_s \rightarrow +\infty$. Donc on a saturation à $+V_{\text{sat}}$.

- c) L'allure de la courbe $v_s(t)$ sur une période est donnée par la figure (S22.5).

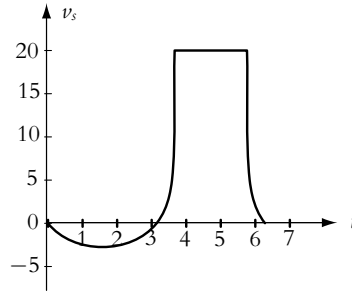


Figure S22.5

3. a) Toujours avec les mêmes conventions, la tension aux bornes de la diode est $v = v_e - v^-$. Or comme précédemment, le fonctionnement linéaire entraîne $v^- = v^+ = 0$, donc $v = v_e$. L'intensité i est égale à $(v^- - v_s)/R$ soit $i = -v_s/R$. En reportant dans l'équation de la caractéristique, on trouve :

$$v_s = -RI_0 \left(\exp \left(\frac{v_e}{V_0} \right) - 1 \right)$$

- b) Pour $t \in [0, \pi]$, la tension v_e est positive donc $v_s < 0$. Le fonctionnement est linéaire tant que :

$$-RI_0 \left(\exp \left(\frac{v_e}{V_0} \right) - 1 \right) > -V_{\text{sat}} \Rightarrow v_e < V_0 \ln \left(1 + \frac{V_{\text{sat}}}{RI_0} \right) \text{ soit } v_e < 25,5 \text{ mV}$$

- c) Pour $t \in [\pi, 0]$, la tension v_e est négative donc $v_s > 0$. Le fonctionnement est linéaire tant que :

$$-RI_0 \left(\exp \left(\frac{v_e}{V_0} \right) - 1 \right) < V_{\text{sat}} \Rightarrow -\exp \left(\frac{v_e}{V_0} \right) + 1 < \frac{V_{\text{sat}}}{RI_0}$$

L'inégalité est vraie pour tout v_e si $V_{\text{sat}} > RI_0$.

4. a) Il s'agit d'un montage sommateur. En appliquant une loi des nœuds à l'entrée inverseuse, sachant qu'avec le fonctionnement linéaire $v^- = v^+ = 0$, on obtient :

$$\frac{v_{s1}}{R_1} + \frac{v_{s2}}{R_1} = -\frac{v_{s3}}{R_1} \Rightarrow v_{s3} = -(v_{s1} + v_{s2})$$

On utilise les relations obtenues pour l'amplificateur logarithmique, ce qui entraîne :

$$v_{s3} = V_0 \left(\ln \left(1 + \frac{v_1}{RI_0} \right) + \ln \left(1 + \frac{v_2}{RI_0} \right) \right)$$

- b) Pour le dernier montage, la relation entre v_s et v_{s3} entraîne :

$$v_s = -RI_0 \left(\exp \left(\frac{v_{s3}}{V_0} \right) - 1 \right) = -RI_0 \left[\left(1 + \frac{v_1}{RI_0} \right) \left(1 + \frac{v_2}{RI_0} \right) - 1 \right]$$

soit $v_s = -v_1 - v_2 - v_1 v_2 / (RI_0)$

- c) Avec les expressions proposées, on peut écrire la valeur moyenne $\langle v_s \rangle$:

$$\langle v_s \rangle = - \langle V_1 \sqrt{2} \cos \omega t \rangle - \langle V_2 \sqrt{2} \cos(\omega t - \varphi) \rangle - \frac{2V_1 V_2}{RI_0} \langle \cos(\omega t) \cos(\omega t - \varphi) \rangle$$

$$\text{Soit } \langle v_s \rangle = - \frac{V_1 V_2}{RI_0} \cos \varphi.$$

- d) La valeur moyenne représente la composante continue du signal, il faut donc disposer d'un filtre passe-bas en sortie du montage de pulsation de coupure inférieure à ω .

- e) L'expression de la puissance est $P = VI \cos \varphi$ sachant que $\cos \varphi = R_{eq} / |Z_{eq}|$ et que $I = V / |Z_{eq}|$,

$$\text{on trouve } P = V^2 \frac{R_{eq}}{R_{eq}^2 + X_{eq}^2}.$$

- f) La résistance r permet d'éviter que v_2 soit nulle ce qui annulerait le produit $V_1 V_2$. Il faut cependant que $r \ll |Z_{eq}|$.

Les lois d'Ohm donnent $V_2 = rI$ et $V_1 = |r + Z_{eq}|I \simeq V_1 = Z_{eq}I$. On obtient alors :

$$\langle v_s \rangle = - \frac{V_1 V_2}{RI_0} \cos \varphi = - \frac{rIZ_{eq}I}{RI_0} \cos \varphi \simeq - \frac{r}{RI_0} IVZ_{eq}$$

donc $k = -RI_0/r$.

Partie V – Mécanique II

Chapitre 23

- A.1** 1. On utilise l'expression générale de la vitesse en coordonnées polaires :

$$\vec{v} = \dot{r}\vec{u}_r + r\dot{\theta}\vec{u}_\theta + \dot{z}\vec{u}_z = R\omega\vec{u}_\theta + h\vec{u}_z$$

2. En calculant la norme de l'expression précédente, on obtient $v = \sqrt{(R\omega)^2 + h^2}$.

3. On utilise l'expression générale de l'accélération en coordonnées polaires :

$$\vec{a} = (\ddot{r} - r\dot{\theta}^2)\vec{u}_r + (2\dot{r}\dot{\theta} + r\ddot{\theta})\vec{u}_\theta + \ddot{z}\vec{u}_z = -R\omega^2\vec{u}_r$$

- B.1** 1. L'expression générale de la vitesse dans la base sphérique est :

$$\vec{v} = \dot{r}\vec{u}_r + r\dot{\theta}\vec{u}_\theta + r\sin\theta\dot{\varphi}\vec{u}_\varphi$$

Le bateau se déplace à la surface de la planète donc $r = R = \text{cte}$. Pour les projections sur les deux autres vecteurs, on se réfère à la figure (S23.1). On se place donc dans un plan localement horizontal, c'est-à-dire tangent à la surface de la planète. Le vecteur $\vec{u}_\theta$ est vers le sud et le vecteur $\vec{u}_\varphi$ est vers l'est.

Le vecteur vitesse fait (en valeur absolue) un angle de $\pi/4$ avec les vecteurs de base et sa projection donne :

$$\vec{v} = -\frac{1}{\sqrt{2}}V_0\vec{u}_\theta + \frac{1}{\sqrt{2}}V_0\vec{u}_\varphi$$

En identifiant avec l'expression générale de $\vec{v}$ on en déduit les deux équations :

$$R\dot{\theta} = -\frac{1}{\sqrt{2}}V_0 \text{ et } R\sin\theta\dot{\varphi} = \frac{1}{\sqrt{2}}V_0$$

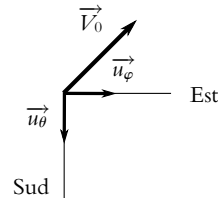


Figure S23.1

2. L'intégration de la première équation donne :

$$\theta = -\frac{V_0}{\sqrt{2}R}t + cte$$

Si l'on suppose qu'à $t = 0$ $\varphi = 0$ et $\theta = \pi/2$ (départ de l'équateur) alors $cte = \pi/2$.
On a alors l'équation différentielle pour φ :

$$\dot{\varphi} = \frac{V_0}{\sqrt{2}R} \frac{1}{\sin\left(-\frac{V_0}{\sqrt{2}R}t + \frac{\pi}{2}\right)} \Rightarrow \varphi = \int \frac{V_0}{\sqrt{2}R} \frac{dt}{\sin\left(-\frac{V_0}{\sqrt{2}R}t + \frac{\pi}{2}\right)} + cte$$

Pour intégrer, il faut effectuer le changement de variable $u = -\frac{V_0}{\sqrt{2}R}t + \frac{\pi}{2}$. On a alors :

$$\varphi = -\int \frac{du}{\sin u} + cte = -\ln\left|\tan \frac{u}{2}\right| + cte$$

À $t = 0$, si $\varphi = 0$ et $u = \theta = \pi/2$ alors $cte = 0$.

3. Le bateau arrive au pôle nord lorsque $\theta = 0$, soit pour $t = \frac{\pi R}{\sqrt{2}V_0}$. On remarquera que φ tend alors vers l'infini, ce qui signifie que le bateau a fait « un nombre infini de tours ». Ce paradoxe est dû au fait qu'on a considéré un point matériel.

Chapitre 24

- A.1** 1. Dans un référentiel galiléen lié au laboratoire, le système est soumis aux forces horizontales $\vec{F}$, $\vec{F}_v$, la tension du ressort $\vec{T}$. Les forces verticales sont le poids et une force de réaction du support $\vec{R}$ (sinon le mouvement ne pourrait être horizontal).

L'application de la relation fondamentale de la dynamique donne :

$$m\vec{a} = m\vec{g} + \vec{F}_v + \vec{T} - k(\ell - \ell_0)\vec{u}_x + \vec{R}$$

avec $\ell = x$. En projection sur l'axe Ox on obtient :

$$m\ddot{x} = -h\dot{x} - k(x - \ell_0) + mF_0 \cos(\omega t)$$

À l'équilibre, il n'y a pas de force d'excitation ; de plus $\ddot{x} = 0$ et $\dot{x} = 0$. La position d'équilibre est donc $x_e = \ell_0$.

Le changement de variable donne, puisque $x_e = cte$: $\ddot{X} = \ddot{x}$, $\dot{X} = \dot{x}$ et finalement :

$$\ddot{X} + \frac{h}{m}\dot{X} + \frac{k}{m}X = F_0 \cos(\omega t)$$

On en déduit $\omega_0^2 = k/m$ et $\tau = 2m/h$.

2. a) On passe en notation complexe pour résoudre l'équation avec : $\underline{X} = X_0 e^{j(\omega t + \varphi)}$. L'équation devient alors :

$$\underline{X} \left(-\omega^2 + \frac{2}{\tau}j\omega + \omega_0^2 \right) = F_0 e^{j\omega t}$$

d'où :

$$X_0(\omega) = \frac{F_0}{\sqrt{(\omega_0^2 - \omega^2)^2 + \left(\frac{2\omega}{\tau}\right)^2}} \text{ et } \varphi = \text{Arg} \frac{F_0}{\omega_0^2 - \omega^2 + \frac{2}{\tau}j\omega}$$

- b) L'expression $X_0(\omega)$ est sous la forme $F_0/\sqrt{f(\omega)}$, donc ses variations sont inverses de celles de $f(\omega)$.
La dérivée de $f(\omega)$ est :

$$\frac{df}{d\omega} = 4\omega \left(\frac{2}{\tau^2} - 2\omega_0^2 + 2\omega^2 \right)$$

On a donc un extremum pour $\omega = 0$ et dans le cas où $\tau > 1/\omega_0$ pour $\omega_r = \omega_0 \sqrt{1 - \frac{1}{(\tau\omega_0)^2}}$. On peut vérifier avec le signe de la dérivée que si ω_r existe c'est minimum pour $f(\omega)$ et un maximum pour $X_0(\omega)$. Dans le cas contraire, la fonction $X_0(\omega)$ est décroissante.

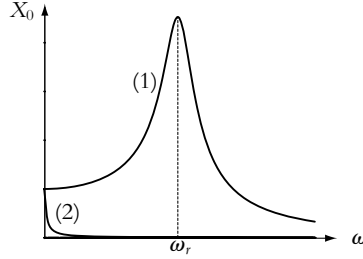


Figure S24.1

En ce qui concerne la phase, comme $F_0 > 0$ alors $\varphi = -\text{Arg} \left(\omega_0^2 - \omega^2 + \frac{2}{\tau}j\omega \right)$. La partie imaginaire est donc négative et $\varphi \in [-\pi, 0]$. De plus $\tan \varphi = -2\omega/(\tau(\omega_0^2 - \omega^2))$. Ainsi :

- Pour $\omega = 0$, $\tan \varphi \rightarrow 0^-$ donc $\varphi \rightarrow 0$.
- Pour $\omega = \omega_0$, $\tan \varphi \rightarrow \pm\infty$ donc $\varphi = -\pi/2$.
- Pour $\omega = \infty$, $\tan \varphi \rightarrow 0^+$ donc $\varphi \rightarrow -\pi$.

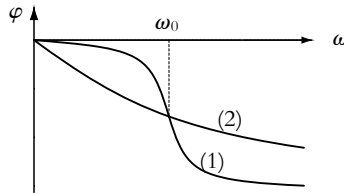


Figure S24.2

Dans les deux figures, le cas (1) correspond à $\tau > 1/\omega_0$ et le cas (2) à $\tau < 1/\omega_0$.

A.2 1. On considère le référentiel du laboratoire galiléen. Les forces s'exerçant sur le mobile sont : le poids, la tension du ressort $\mathcal{T}$, la force de frottement et la force d'excitation. On applique la relation fondamentale de la dynamique en projection sur l'axe Oz descendant :

$$m\ddot{z} = -k(z - \ell_0) - h\dot{z} + mg + mF_0 \cos \omega t$$

Pour avoir l'équation en $\dot{z}$, on dérive :

$$m\ddot{v} = -kv - h\dot{v} - mF_0\omega \sin \omega t \Rightarrow \ddot{v} + \frac{\omega_0}{Q}\dot{v} + \omega_0^2 v = -F_0\omega \sin \omega t$$

2. a) En notation complexe, le second membre de l'équation s'écrit $F_0\omega e^{j(\omega t + \pi/2)}$ ou $F_0j\omega e^{j\omega t}$. L'équation différentielle devient donc :

$$\mathcal{L} \left(-\omega^2 + j\frac{\omega\omega_0}{Q} + \omega_0^2 \right) = j\omega F_0 e^{j\omega t} \Rightarrow \mathcal{L} \left(j\omega + \frac{\omega_0}{Q} - j\frac{\omega_0^2}{\omega} \right) = F_0 e^{j\omega t}$$

On en déduit le module et la phase :

$$V_0(\omega) = \frac{F_0}{\sqrt{\frac{\omega_0^2}{Q^2} + \omega_0^2 \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)^2}} \text{ et } \varphi = \text{Arg} \left(\frac{F_0}{\frac{\omega_0}{Q} + j\omega_0 \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)} \right)$$

- b) Pour l'étude de $V_0(\omega)$, on remarque que le numérateur est constant et que sous la racine du dénominateur, on a la somme d'un terme constant et d'une fonction positive ou nulle puisque c'est un carré. Le dénominateur est donc minimal lorsque cette fonction est nulle soit pour $\omega = \omega_0$, ce qui correspond à un maximum pour $V_0(\omega)$ égal à $F_0 Q / \omega_0$. Pour $\omega = 0$ et $\omega \rightarrow \infty$, $V_0(\omega) \rightarrow 0$. En ce qui concerne la phase, l'argument du numérateur positif est égal à 0 et si l'on note φ_d l'argument du dénominateur, on a donc $\varphi = -\varphi_d$. La partie réelle du dénominateur est positive, donc $\cos \varphi_d > 0$ et $\varphi_d \in \left[-\frac{\pi}{2}, \frac{\pi}{2}\right]$. De plus on a

$$\tan \varphi_d = \frac{\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega}}{Q}.$$

$$\text{Ainsi } \varphi = -\arctan \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right) / Q.$$

On a donc :

- Pour $\omega = 0$, $\tan \varphi_d \rightarrow -\infty$ donc $\varphi_d \rightarrow -\pi/2$ et $\varphi \rightarrow \pi/2$.
- Pour $\omega = \omega_0$, $\tan \varphi_d = 0$ donc $\varphi = 0$.
- Pour $\omega \rightarrow +\infty$, $\tan \varphi \rightarrow +\infty$ donc $\varphi_d \rightarrow \pi/2$ et $\varphi \rightarrow -\pi/2$.

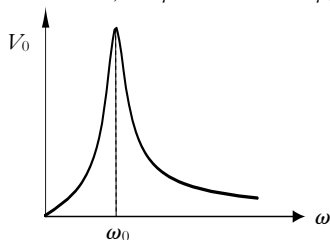


Figure S24.3 Amplitude

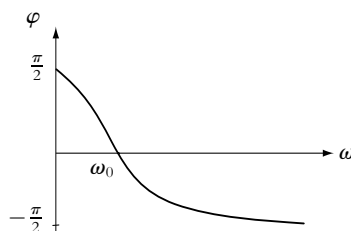


Figure S24.4 Phase

3. a) La puissance cinétique est $P_c = \frac{dE_c}{dt}$ soit avec $v = V_0 \cos(\omega t + \varphi)$:

$$P_c = -mV_0^2 \omega \cos(\omega t + \varphi) \sin(\omega t + \varphi) = \frac{-mV_0^2}{2} \sin(2\omega t + \varphi) \Rightarrow \langle P_c \rangle = 0$$

Pour la puissance du ressort P_r , il suffit de dériver l'énergie potentielle. Si l'on note Z l'élongation du ressort, l'énergie potentielle est $E_p = kZ^2/2$. On obtient Z en intégrant v , d'où $X = V_0/\omega \sin(\omega t + \varphi)$ et :

$$P_r = \frac{dE_p}{dt} = \frac{kV_0^2}{\omega} \sin(\omega t + \varphi) \cos(\omega t + \varphi) = \frac{kV_0^2}{2\omega} \sin(2\omega t + \varphi) \Rightarrow \langle P_r \rangle = 0$$

- b) Si l'on reprend l'équation du mouvement que l'on multiplie par v :

$$m\ddot{z}v = F_r v + F_r v + mgv + Fv \Rightarrow \frac{dE_c}{dt} = P_r + P_f + P_{mg} + P_e$$

où F_r est la force du ressort. On peut donc écrire

$$\langle P_c \rangle = \langle P_r \rangle + \langle P_f \rangle + \langle P_{mg} \rangle + \langle P_e \rangle$$

La puissance moyenne du poids est nulle car $\langle P_{mg} \rangle = mg \langle v \rangle$. On trouve donc $\langle P_f \rangle + \langle P_e \rangle = 0$.

D'autre part $\langle P_f \rangle = -h \langle v^2 \rangle = -hV_0^2/2$. Cette puissance moyenne est donc maximale lorsque $V_0(\omega)$ est maximale donc à la résonance.

- A.3** 1. On écrit la relation fondamentale de la dynamique dans le référentiel galiléen :

$$m \frac{d\vec{OM}}{dt} = -k\vec{OM} - \lambda \frac{d\vec{OM}}{dt} - eE_0 \cos \omega t \vec{u}_x$$

2. La solution forcée par la force électrique aura même direction que cette force, elle est donc suivant $\vec{u}_x$. On projette donc sur l'axe Ox , soit :

$$m\ddot{X} + \lambda\dot{X} + kX = -eE_0 \cos \omega t \Rightarrow \ddot{X} + \frac{\omega_0}{Q}\dot{X} + \omega_0^2 X = -\frac{eE_0}{m} \cos \omega t$$

avec $\omega_0 = \sqrt{k/m}$ et $Q = m\omega_0/\lambda$.

On cherche une solution de la forme $X = X_0 \cos(\omega t + \varphi)$ et on utilise les complexes, ce qui entraîne :

$$\underline{X} \left(-\omega^2 + j\frac{\omega\omega_0}{Q} + \omega_0^2 \right) = \frac{-eE_0}{m} e^{j\omega t} \Rightarrow X_0 e^{j\varphi} = \frac{-eE_0}{m} \frac{1}{\omega_0^2 - \omega^2 + j\frac{\omega\omega_0}{Q}}$$

3. Les applications numériques donnent $\omega_0 = 1,05 \cdot 10^{16} \text{ rad.s}^{-1}$ et $Q = 9,6 \cdot 10^5$.
 4. Pour le rouge (800 nm), on trouve $\omega_r = 2,3 \cdot 10^{15} \text{ rad.s}^{-1}$ et pour le bleu (450 nm), $\omega_b = 4,2 \cdot 10^{15} \text{ rad.s}^{-1}$. On calcule alors que le rapport ω^2/ω_0^2 est dans l'intervalle $[0,05; 0,2]$. Étant donné la valeur importante de Q , on peut négliger le terme $\omega\omega_0/Q$ devant les autres. On écrira donc :

$$X_0 = \frac{eE_0}{m} \frac{1}{|\omega_0^2 - \omega^2|}$$

5. L'accélération est $\ddot{X}$ donc proportionnelle à $\omega^2 X_0$. La puissance réémise est donc proportionnelle à $\omega^4 X_0^2$. On en déduit le rapport de la puissance émise dans le bleu sur la puissance émise dans le rouge :

$$\frac{P_b}{P_r} = \left(\frac{\omega_b}{\omega_r} \right)^4 \left(\frac{\omega_b^2 - \omega_0^2}{\omega_r^2 - \omega_0^2} \right)^2 \simeq 16$$

Les molécules du ciel réémettent beaucoup plus dans le bleu d'où la couleur bleue du ciel.

- B.1** 1. Par symétrie du système, et comme les points ne subissent pas d'autres forces horizontales que les tensions des ressorts, on peut considérer que les ressorts (1) et (3) ont même longueur à l'équilibre. On écrit les relations d'équilibres pour chaque point :

Pour A : $0 = -k(\ell_{eq} - \ell_0) + K(L_{eq} - L_0)$ et pour B on obtient la même équation.

2. On note ℓ_1 , ℓ_2 et ℓ_3 les longueurs des ressorts. La relation fondamentale de la dynamique pour A donne :

$$m\ddot{x} = -k(\ell_1 - \ell_0) + K(\ell_2 - L_0) + mF_0 \cos \omega t$$

or $\ell_1 = \ell_{eq} + x$ et $\ell_2 = L_{eq} - x + y$. En reportant ces grandeurs dans l'équation précédente et en retranchant l'équation d'équilibre, on trouve :

$$\ddot{x} + \frac{k+K}{m}x = \frac{K}{m}y + F_0 \cos \omega t$$

La relation fondamentale de la dynamique pour B donne :

$$m\ddot{y} = k(\ell_3 - \ell_0) - K(\ell_2 - L_0)$$

or $\ell_3 = \ell_{eq} - y$. En reportant dans l'équation précédente et en retranchant l'équation d'équilibre, on trouve :

$$\ddot{y} + \frac{k+K}{m}y = \frac{K}{m}x$$

3. a) On écrit les équations précédentes sous forme complexe, ce qui donne le système suivant :

$$\begin{cases} \underline{X} \left(-\omega^2 + \frac{k+K}{m} \right) - \frac{K}{m} \underline{Y} = F_0 e^{i\omega t} \\ \underline{Y} \left(-\omega^2 + \frac{k+K}{m} \right) - \frac{K}{m} \underline{X} = 0 \end{cases}$$

En résolvant le système on trouve les solutions de l'énoncé. Les différentes constantes sont : $\omega_0^2 = k/m$, $\omega_1^2 = (k+2K)/m$, $\omega_2^2 = (k+K)/m$ et $\omega_3^2 = K/m$.

- b) On en déduit X_0 et Y_0 en prenant les modules des expressions précédentes :

$$X_0 = \frac{|\omega_2^2 - \omega^2|}{|\omega_0^2 - \omega^2||\omega_1^2 - \omega^2|} F_0 \text{ et } Y_0 = \frac{\omega_3^2}{|\omega_0^2 - \omega^2||\omega_1^2 - \omega^2|} F_0$$

Quant aux phases, elles valent 0 ou π suivant le signe des rapports. Sachant que $\omega_0 < \omega_2 < \omega_1$, on a :

- Si $\omega < \omega_0$ alors $\varphi = \psi = 0$.
- si $\omega_0 < \omega < \omega_2$ alors $\varphi = \psi = \pi$.
- si $\omega_2 < \omega < \omega_1$ alors $\varphi = 0$ et $\psi = \pi$.
- Si $\omega_1 < \omega$ alors $\varphi = \pi$ et $\psi = 0$.

- c) Les courbes de X_0 et Y_0 sont données sur les figures (S24.5) et (S24.6)

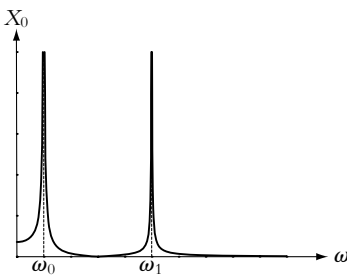


Figure S24.5

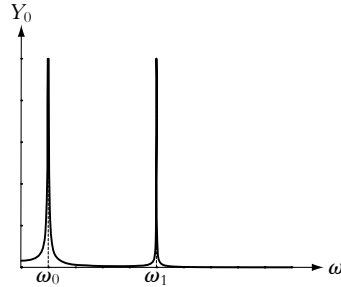


Figure S24.6

Les résonances infinies sont dues au fait qu'on n'a pas tenu compte des frottements.

- B.2** 1. Une force de frottement fluide est une force du type $\vec{F}_v = -D\vec{v}$, D étant une constante positive et $\vec{v}$ la vitesse du mobile.
2. a) La relation de Chasles donne $\vec{OM} = \vec{OA} + \vec{AB} + \vec{BM}$ sachant que $\vec{AB}$ est une constante. À l'équilibre, il n'y a que le poids et la force du ressort :

$$0 = mg - k(AB + \xi_{\text{eq}} - \ell_0) \Rightarrow 0 = mg - k(AB - \ell_0)$$

puisque $\xi_{\text{eq}} = 0$. On écrit la relation fondamentale de la dynamique dans le référentiel galiléen. Le système d'amortissement étant fixé à A c'est la vitesse de M par rapport à A qui va intervenir :

$$M \frac{d\vec{OM}}{dt} = -k(\vec{AM} - \ell_0 \vec{u}_z) + m\vec{g} - D \frac{d\vec{AM}}{dt}$$

On développe et on projette sur Oz en utilisant l'équation d'équilibre pour simplifier :

$$M(\ddot{u} + \ddot{\xi}) = -k\xi - D\dot{\xi} \Rightarrow \ddot{\xi} + \frac{D}{M}\dot{\xi} + \frac{k}{M} = -\ddot{u}$$

ce qui est l'équation demandée.

- b) Si en régime forcé ξ est de la forme $\xi_0 \cos(\omega t + \varphi)$ alors l'ordre de grandeur de l'amplitude de $\dot{\xi}$ est $\omega \xi_0$ et l'ordre de grandeur de l'amplitude de $\ddot{\xi}$ est $\omega^2 \xi_0$, donc aux grandes fréquences c'est le terme $\ddot{\xi}$ qui est prédominant. L'équation se réduit donc à $\ddot{\xi} = -\ddot{u}$. En intégrant, on a donc $\dot{\xi} = -\dot{u}$ et ξ représente le mouvement du sol par rapport à $\mathcal{R}$: on a fabriqué un sismographe.
- c) Dans le cas des mouvements lents, c'est le terme en ξ qui est prédominant soit $\xi = -\frac{\ddot{u}}{\omega_s^2}$ et ξ mesure l'accélération.
3. a) On cherche ξ de la forme $\xi_0 \cos(\omega t + \varphi)$ et on passe en notation complexe soit $\underline{\xi} = X e^{i(\omega t + \varphi)}$. L'équation différentielle devient :

$$\underline{\xi}(-\omega^2 + 2\eta i\omega + \omega_s^2) = -U_0 \omega^2 e^{i\omega t} \Rightarrow \underline{\xi} = \frac{-U_0 \omega^2 e^{i\omega t}}{(-\omega^2 + 2\eta i\omega + \omega_s^2)}$$

- b) En introduisant les notations de l'énoncé on peut calculer le module de ξ :

$$X(\omega) = \frac{U_0 x^2}{\sqrt{(1-x^2)^2 + 4h^2 x^2}}$$

Le déphasage Φ est égal à celui du numérateur (π car il y a un signe "-") moins celui du dénominateur (φ_d) d'où :

$$\Phi = \pi - \varphi_d = \pi - \text{Arg}(1 - x^2 + i2hx)$$

- c) Les comportements limites pour X sont les suivants :

- Lorsque $\omega \rightarrow 0$ alors $x \rightarrow 0$ et $X \sim U_0 x^2$ donc $X \rightarrow 0$.

- Lorsque $\omega \rightarrow +\infty$ alors $x \rightarrow +\infty$ et $X \sim U_0 x^2 / \sqrt{x^4}$ donc $X \rightarrow U_0$.

En ce qui concerne la phase, la partie imaginaire du dénominateur de ξ est positive donc $\sin \varphi_d > 0$ et $\varphi_d \in [0, \pi]$. On peut déterminer $\tan \varphi_d = \frac{2hx}{1-x^2}$. Les comportements limites pour φ sont les suivants :

- Si $x \rightarrow 0$, $\tan \varphi_d \rightarrow 0^+$ alors $\varphi_d \rightarrow 0$ et $\Phi \rightarrow \pi$.

- Si $x \rightarrow +\infty$, $\tan \varphi_d \rightarrow 0^-$ alors $\varphi_d \rightarrow \pi$ et $\Phi \rightarrow 0$.

- d) Le changement de variable donne $Y = (x'^2 - 1)^2 + 4h^2 x'^2$. La dérivée de cette fonction par rapport à x' est :

$$\frac{dY}{dx'} = 4x'(2h^2 + x'^2 - 1)$$

Pour que cette fonction ait un extremum autre que $x' = 0$, il faut que $h < 1/\sqrt{2}$ et alors l'extremum qui correspond à un minimum pour Y et un maximum pour X (résonance) a lieu pour $x' = \sqrt{1 - 2h^2}$.

- e) Les courbes pour X (figure S24.7) et pour Φ (figure S24.8) correspondent à $h = 1$ pour les courbes (1) sans résonance et $h = 0,5$ pour les courbes (2) avec résonance. Pour le tracé on a pris $U_0 = 1$.

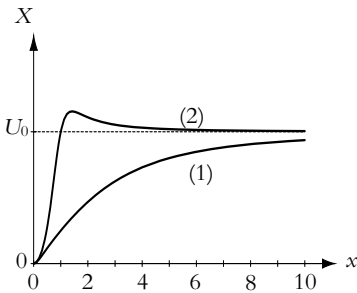


Figure S24.7 Amplitude

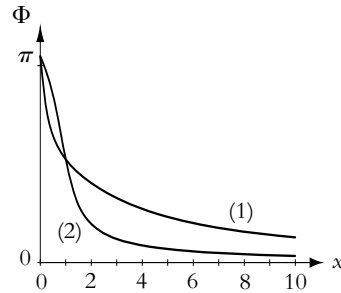


Figure S24.8 Phase

4. a) Pour les mouvements lents (petites pulsations), on a vu que $X \sim U_0 x^2$ ou $X \sim U_0 \frac{\omega^2}{\omega_s^2}$. La variation d'accélération Δa est égale à $\omega^2 U_0$ d'où $\sigma = 1/\omega_s$.
 b) Sachant que $\omega_s = 2\pi/T_0$, l'application numérique donne $\Delta \xi = 0,1 \text{ mm}$.

B.3 1. Dans le référentiel galiléen lié au sol, le système M subit à l'équilibre, le poids et la tension du ressort $\vec{T} = -k(\ell - \ell_0)\vec{u}_z$. La longueur ℓ du ressort est égale à $z_G - R$. La relation de la statique à l'équilibre projetée sur $\vec{u}_z$ donne :

$$0 = -Mg - k(\ell_e - \ell_0) \Rightarrow z_{G,\text{eq}} = -\frac{Mg}{k} + \ell_0 + R$$

2. a) On retrouve l'énergie potentielle E_p de pesanteur en écrivant le travail élémentaire du poids :

$$dE_{pg} = -\delta W = -Md dz_G \Rightarrow E_{pg} = Mgz_G + cte$$

- b) Même méthode pour l'énergie potentielle élastique :

$$dE_{pr} = -\delta W = -k(\ell - \ell_0) d\ell = -k(\ell - \ell_0) d(\ell - \ell_0) \Rightarrow E_{pr} = \frac{k}{2}(\ell - \ell_0)^2 + cte$$

soit avec z_G , $E_{pr} = \frac{k}{2}(z_G - R - \ell_0)^2 + cte$

- c) Dans le cas du mouvement il faut ajouter aux forces précédentes, la force de frottement qui vaut $\vec{F} = -\lambda \dot{z}_G \vec{u}_z$ puisque la route est plane. Le théorème de la puissance cinétique s'écrit :

$$\frac{dE_c}{dt} = \mathcal{P}(m\vec{g}) + \mathcal{P}(\vec{T}) + \mathcal{P}(\vec{F}) = -\frac{dE_{pg}}{dt} - \frac{dE_{pr}}{dt} - \lambda \dot{z}_G^2$$

Le calcul des dérivées donne :

$$\dot{z}_G(M\ddot{z}_G + Mg + k(z_G - R - \ell_0) + \lambda \dot{z}_G) = 0$$

On exclut l'immobilité donc $\dot{z}_G = 0 \forall t$ et avec l'équation à l'équilibre (en utilisant $z = z_G - z_{G,eq}$), on trouve finalement :

$$M\ddot{z} + \lambda \dot{z} + kz = 0$$

- d) On met l'équation différentielle sous la forme :

$$\ddot{z} + \frac{\lambda}{M}\dot{z} + \frac{k}{M}z = 0$$

On calcule le discriminant de l'équation caractéristique :

$$\Delta = \left(\frac{\lambda}{M}\right)^2 - 4\frac{k}{M} = -384 < 0$$

Le régime est pseudopériodique.

- e) La relation fondamentale de la dynamique donne :

$$M\vec{a} = M\vec{g} + \vec{T} + \vec{F} \Rightarrow M\ddot{z}_G = Mg - k(z_G - R - \ell_0) - \lambda \dot{z}_G$$

- f) La résolution classique de l'équation caractéristique donne comme racines $\frac{-\lambda}{2M} \pm i\frac{\sqrt{-\Delta}}{2}$ soit

$$z(t) = A \exp\left(-\frac{\lambda t}{2M}\right) \cos(\Omega t + \varphi) \text{ avec } \Omega = \frac{\sqrt{-\Delta}}{2}$$

L'allure de la courbe est représentée sur la figure (S24.9) sachant qu'à $t = 0$, $z = 0$ tandis que $\dot{z} \neq 0$ en raison de l'impulsion soudaine.

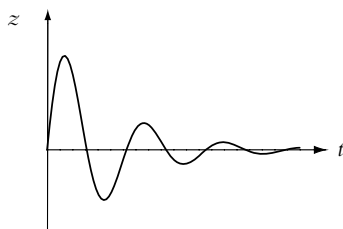


Figure S24.9

3. L'unité de λ est $N.m^{-1}.s$ ou $kg.s^{-1}$.
4. La période spatiale est égale à la vitesse multipliée par la période temporelle (comme une longueur d'onde) soit $L = vT$ or $\omega = 2\pi/T = 2\pi v/L$.
5. La relation fondamentale de la dynamique en projection sur Oz donne :

$$M\ddot{z}_G = -Mg - \lambda(\dot{z}_G - \dot{z}_0) - k(\ell - \ell_0)$$

Sachant que $\ell = z_G - z_0$ et que $z = z_G - z_{G,eq}$, avec $z_{G,eq}$ déterminé dans la première partie, on a $\dot{z} = \dot{z}_G$ et $\ddot{z} = \ddot{z}_G$. L'équation précédente se simplifie pour donner :

$$M\ddot{z} = -kz + kz_0 - kR + \lambda \dot{z}_0$$

6. La solution de l'équation différentielle est la somme de la solution de l'équation homogène (régime transitoire) et la solution particulière imposée par le second membre (régime forcé). Le régime transitoire s'atténue rapidement et on ne s'intéresse qu'au régime forcé qui sera imposé par la forme de z_0 et $\dot{z}_0$, donc régime sinusoïdal de pulsation ω .

7. a) On introduit les grandeurs complexes dans l'équation précédente :

$$\underline{Z}(-M\omega^2 + \lambda j\omega + k) = \underline{A}(k + \lambda j\omega) \Rightarrow \frac{\underline{Z}}{\underline{A}} = \frac{k + \lambda j\omega}{-M\omega^2 + \lambda j\omega + k}$$

En divisant numérateur et dénominateur par M puis en mettant k/M en facteur au numérateur et au dénominateur, on trouve la forme demandée avec $\omega_1 = k/\lambda$, $\omega_0 = \sqrt{k/M}$ et $Q = \sqrt{kM}/\lambda$.

- b) L'application numérique donne : $\omega_0 = 10 \text{ rad.s}^{-1}$, $\omega_1 = 25 \text{ rad.s}^{-1}$ et $Q = 2, 5$.

c)

$$\left| \frac{\underline{Z}}{\underline{A}} \right| = \sqrt{\frac{1 + \frac{\omega^2}{\omega_1^2}}{\left(1 - \frac{\omega^2}{\omega_0^2}\right)^2 + \left(\frac{\omega}{Q\omega_0}\right)^2}}$$

8. a) On étudie les fonctions :

$$|\underline{H}_1| = \sqrt{1 + \frac{\omega^2}{\omega_1^2}} \text{ et } |\underline{H}_2| = \frac{1}{\sqrt{\left(1 - \frac{\omega^2}{\omega_0^2}\right)^2 + \left(\frac{\omega}{Q\omega_0}\right)^2}}$$

soit pour le gain en dB :

$$G_1 = 20 \log \left(1 + \frac{\omega^2}{\omega_1^2}\right)^{1/2} \text{ et } G_2 = -20 \log \left(\left(1 - \frac{\omega^2}{\omega_0^2}\right)^2 + \left(\frac{\omega}{Q\omega_0}\right)^2\right)^{1/2}$$

Étude asymptotique de G_1 :

- Pour $\omega \ll \omega_1$, $G_1 \sim 0$: asymptote horizontale.
- Pour $\omega \gg \omega_1$, $G_1 \sim 20 \log \frac{\omega}{\omega_1}$: asymptote de pente +20 dB/dec.

Étude asymptotique de G_2 :

- Pour $\omega \ll \omega_0$, $G_2 \sim 0$: asymptote horizontale.
- Pour $\omega \gg \omega_0$, $G_2 \sim -40 \log \frac{\omega}{\omega_0}$: asymptote de pente -40 dB/dec.

Sur la figure (S24.10), on a représenté les diagrammes asymptotiques de G_1 , G_2 et G ainsi que la courbe G .

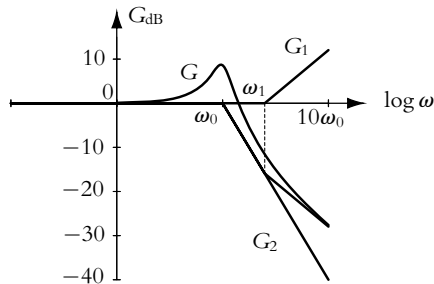


Figure S24.10

- b) La valeur de ν correspondante est $1,59 \text{ m.s}^{-1}$ soit $5,7 \text{ km.h}^{-1}$. Pour $\omega = \omega_0$; on a :

$$\frac{Z}{A} = Q \sqrt{1 + \frac{\omega_0^2}{\omega_1^2}} = 2,69$$

L'amplitude des oscillations du véhicule est alors 26,9 cm ce qui est très important.

9. Pour éviter des oscillations trop importantes, il faut éviter la zone de résonance.

B.4 1. On étudie le système constitué de la masse M dans le référentiel terrestre considéré comme galiléen. Ce système est soumis à son poids et à la force de rappel du ressort. La projection du principe fondamental de la dynamique à l'équilibre sur l'axe Oz s'écrit $Mg - k(h - l_0) = 0$.

2. En faisant la même chose hors d'équilibre, on obtient $\ddot{z} + \frac{k}{M}z = \frac{kh}{M}$.

3. On a donc des oscillations harmoniques de pulsation $\omega_0 = \sqrt{\frac{k}{M}}$ (avec les données fournies à ce niveau, il n'est pas possible de donner l'amplitude).

4. Lorsque la masse est dans l'eau, on ajoute la poussée d'Archimède dans le bilan des forces. Il suffit de remplacer le poids $M\vec{g}$ par $M\vec{g} \left(1 - \frac{\rho_{\text{eau}}}{\rho}\right)$ dans ce qui précède. La forme de l'équation n'est pas

modifiée mais la position de l'équilibre devient $\frac{Mg \left(1 - \frac{\rho_{\text{eau}}}{\rho}\right) + kl_0}{k}$. On notera que la pulsation n'est pas modifiée.

5. On ajoute maintenant la force $-\alpha\vec{v}$ et on en déduit l'équation du mouvement suivante $\ddot{z} + \frac{\alpha}{M}\dot{z} + \frac{k}{M}z = \frac{kh}{M}$.

6. Lorsque l'amortissement est faible, α prend une faible valeur. Dans ce cas, le discriminant de l'équation caractéristique $\Delta = \left(\frac{\alpha}{M}\right)^2 - 4\frac{k}{M}$ est négatif. Le régime est pseudo-périodique avec à l'instant initial une valeur h_1 , une tangente horizontale due à une vitesse initiale nulle et une valeur finale à h .

7. On ajoute enfin une force excitatrice soit l'équation du mouvement toujours obtenue par projection du principe fondamental de la dynamique sur Oz $\ddot{z} + \frac{\alpha}{M}\dot{z} + \frac{k}{M}z = \frac{kh}{M} + \omega^2 a \cos \omega t$.

8. On passe à la notation complexe et on en déduit $(-M\omega^2 + j\alpha\omega + k)Z e^{j\varphi} = m\omega^2 A$ soit une ampli-

$$\text{tude } Z = \frac{x^2 a}{\sqrt{(1 - x^2)^2 + \frac{x^2}{\tau^2 \omega_0^2}}}.$$

9. L'amplitude Z tend vers 0 pour des pulsations très faibles et vers a pour des pulsations très grandes. L'amplitude Z est supérieur à a si on a un passage par un maximum : on recherche donc un extremum. Pour cela, on calcule la dérivée de la fonction $f(X) = (1 - X)^2 + \frac{X}{\omega_0^2 \tau^2}$ soit

$$f'(X) = 2X + \left(\frac{1}{\omega_0^2 \tau^2} - 2\right). \text{ Par conséquent, on a un extremum uniquement si } M > \frac{\alpha^2}{2k}.$$

Dans ce cas, l'amplitude Z vérifie $Z > a$ pour $\omega > \frac{\omega_0}{\sqrt{2 - \frac{1}{\omega_0^2 \tau^2}}}$.

10. Le maximum est obtenu pour $\omega_r = \frac{\omega_0}{\sqrt{1 - \frac{1}{2\omega_0^2 \tau^2}}}$.

11. La valeur maximale est obtenue en reportant la valeur de ω_r dans l'expression de l'amplitude soit $Z_{\text{max}} = \frac{2aMk}{\alpha\sqrt{4Mk - \alpha^2}}$.

Chapitre 25

- A.1** 1. Le référentiel lié au sol est considéré galiléen. Les forces s'exerçant sur le système M sont : le poids et la réaction normale $\vec{R}_N$ de la route. Les forces sont représentées sur le schéma de la figure (S25.1).

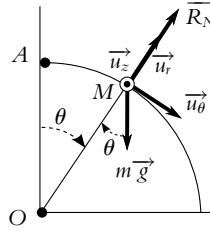


Figure S25.1

Le moment cinétique du point M par rapport à O est $\vec{\sigma}_0 = \vec{OM} \wedge m\vec{v}$. Avec $\vec{OM} = R\vec{u}_r$ et puisque le mouvement est circulaire $\vec{v} = R\dot{\theta}\vec{u}_\theta$, le moment cinétique est $\vec{\sigma}_0 = mR^2\dot{\theta}\vec{u}_z$. Le théorème du moment cinétique par rapport à O s'écrit :

$$\frac{d\vec{\sigma}_0}{dt} = \vec{OM} \wedge m\vec{g} + \vec{OM} \wedge \vec{R}_N = \vec{OM} \wedge m\vec{g}$$

Le produit vectoriel $\vec{OM} \wedge \vec{R}_N$ est nul puisque $\vec{R}_N$ est colinéaire à $\vec{OM}$. Le calcul donne :

$$mR^2\ddot{\theta} = mgR \sin \theta \Rightarrow \ddot{\theta} = \frac{g}{R} \sin \theta$$

2. On multiplie l'équation précédente par $\dot{\theta}$ ce qui permet de l'intégrer :

$$\dot{\theta}\ddot{\theta} = \frac{g}{R}\dot{\theta} \sin \theta \Rightarrow \frac{\dot{\theta}^2}{2} = -\frac{g}{R} \cos \theta + cte$$

À $t = 0$, la vitesse du point est v_0 soit $\dot{\theta}(0) = v_0/R$ et $\theta(0) = 0$, ce qui permet de déterminer la constante :

$$\frac{v_0^2}{2R^2} = -\frac{g}{R} + cte \Rightarrow \frac{\dot{\theta}^2}{2} - \frac{v_0^2}{2R^2} = -\frac{g}{R}(\cos \theta - 1)$$

La vitesse angulaire va augmenter au cours du mouvement car il n'y a pas de forces de frottement donc $\dot{\theta} > 0$ et :

$$\dot{\theta} = \sqrt{\frac{2g}{R}(1 - \cos \theta) + \frac{v_0^2}{R^2}}$$

Deux remarques :

- Le fait de multiplier par $\dot{\theta}$ pour intégrer l'équation doit faire penser à la simplification par ce même terme $\dot{\theta}$ lorsqu'on dérive l'énergie mécanique pour trouver l'équation du mouvement. On a réalisé ici la procédure inverse.
- Comme on intègre par rapport au temps, il est nécessaire de multiplier par $\dot{\theta}$ avant d'intégrer.

- A.2** 1. La force électrostatique exercée par le proton (en O) sur l'électron (en M) est :

$$\vec{f} = \frac{-e^2}{4\pi\epsilon_0 R^2} \vec{u}_r \text{ avec } \vec{u}_r = \frac{\vec{OM}}{R}$$

L'énergie potentielle dont dérive la force $\vec{F}$ est $E_p = \frac{-e^2}{4\pi\epsilon_0 R}$. On peut alors écrire l'énergie totale égale à l'énergie mécanique et en déduire une expression de v :

$$E_m = \frac{1}{2}mv^2 - \frac{e^2}{4\pi\epsilon_0 R} = -\frac{k}{n^2} \Rightarrow v^2 = \frac{2}{m} \left(-\frac{k}{n^2} + \frac{e^2}{4\pi\epsilon_0 R} \right)$$

On demande une expression de ν indépendante de R , il faut donc trouver une équation supplémentaire : par exemple la relation fondamentale de la dynamique en coordonnées polaires :

$$m \vec{a} = \vec{f} \Rightarrow -mR\dot{\theta}^2 = \frac{-e^2}{4\pi\epsilon_0 R^2} \Rightarrow v^2 = (R\dot{\theta})^2 = \frac{e^2}{4m\pi\epsilon_0 R}$$

En reportant dans l'équation précédente, on en déduit $\nu = \sqrt{\frac{2k}{m}} \frac{1}{n}$ et avec la relation entre R et ν^2 :

$$R = \frac{e^2}{8\pi k \epsilon_0} n^2.$$

2. Il s'agit d'un mouvement circulaire donc le moment cinétique par rapport au centre s'écrit :

$$\vec{\sigma}_0 = \overrightarrow{OM} \wedge m \vec{v} \Rightarrow \vec{\sigma}_0 = mR^2 \dot{\theta} \vec{u}_z = mR\nu \vec{u}_z = n \sqrt{\frac{2m}{k}} \frac{e^2}{8\pi\epsilon_0}$$

- A.3** 1. Les forces s'exerçant sur le point M sont :

- le poids $m \vec{g} = -mg \vec{u}_z$,
- la réaction normale $\vec{R}_N = R_N \vec{u}_z$,
- la force de tension du fil $\vec{T} = -F \vec{u}_r$.

Le poids et la réaction normale se compensent. Il reste donc la tension qui est une force centrale de centre O . On en déduit que le moment cinétique par rapport à O se conserve.

On détermine le moment cinétique en utilisant les coordonnées cylindriques :

$$\vec{\sigma}_0 = \overrightarrow{OM} \wedge m \vec{v} \Rightarrow \vec{\sigma}_0 = r \vec{u}_r \wedge m(\dot{r} \vec{u}_r + r \dot{\theta} \vec{u}_\theta) = mr^2 \dot{\theta} \vec{u}_z$$

La vitesse aréolaire est la vitesse à laquelle le rayon vecteur $\overrightarrow{OM}$ « balaie » une surface unité. On se reportera au cours sur les forces centrales pour montrer que la norme de cette vitesse est égale à $C/2$ où C est la constante des aires égale à $|\vec{\sigma}_0|/m = r^2 \dot{\theta}$.

Le mouvement est à force centrale donc il suit la loi des aires : le rayon vecteur balaie des surfaces égales en des temps égaux.

2. Le force de traction (tension du ressort) est $\vec{T} = -F \vec{u}_r$ avec $F = cte$. On peut calculer son énergie potentielle en calculant le travail :

$$\delta W = \vec{T} \cdot d\vec{r} = -F \vec{u}_r \cdot (dr \vec{u}_r + r d\theta \vec{u}_\theta) = -F dr$$

En intégrant avec $\delta W = -dE_p = -F dr$, on en déduit $E_p = Fr + cte$. La constante est nulle puisque on prend $E_p = 0$ pour $r = 0$.

Les deux autres forces sont perpendiculaires au déplacement donc elles ne travaillent pas. L'énergie mécanique est donc :

$$E_m = E_c + E_p = \frac{1}{2} m(\dot{r}^2 + r^2 \dot{\theta}^2) + Fr$$

3. a) Dans l'expression de l'énergie, on remplace $\dot{\theta}$ par C/r^2 ce qui donne :

$$E_m = \frac{1}{2} m \dot{r}^2 + \frac{mC^2}{2r^2} + Fr = \frac{1}{2} m \dot{r}^2 + E_{p,\text{eff}} \text{ avec } E_{p,\text{eff}} = \frac{mC^2}{2r^2} + Fr$$

où $E_{p,\text{eff}}$ est l'énergie potentielle efficace ou effective. On obtient donc une équation du mouvement qui ne dépend explicitement que de r et de $\dot{r}$: on s'est donc ramené à un problème à une dimension le long de la droite OM sachant que la vitesse angulaire est obtenue par $\dot{\theta} = C/r^2$.

- b) On peut déterminer l'expression de la constante des aires en fonction des données initiales. À $t = 0$, la vitesse V_0 est perpendiculaire à $\overrightarrow{OM}_0 = D \vec{u}_r$, donc $\vec{\sigma}_0 = mDV_0 \vec{u}_z$ et $C = DV_0$. On calcule la dérivée de $E_{p,\text{eff}}$ par rapport à r :

$$\frac{dE_{p,\text{eff}}}{dr} = \frac{-mC^2}{r^3} + F \Rightarrow \frac{dE_{p,\text{eff}}}{dr} = 0 \text{ pour } r_1 = \left(\frac{mC^2}{F} \right)^{1/3} = \left(\frac{mD^2 V_0^2}{F} \right)^{1/3}$$

Pour $r < r_1$, la dérivée est négative et pour $r > r_1$ elle est positive donc r_1 correspond à un minimum.

- c) On a représenté l'énergie potentielle efficace sur la figure (S25.2). La droite horizontale représente l'énergie mécanique qui est constante. Le mouvement ne peut avoir lieu que dans la zone $E_m \geq E_{p,\text{eff}}$ puisque $m\dot{r}^2/2 > 0$. Cette zone est limitée par deux valeurs de $r = OM$ notées r_m et r_M sur la figure.

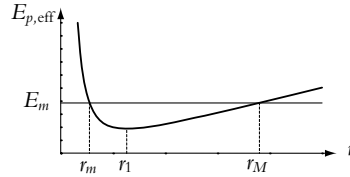


Figure S25.2

La position initiale correspond à r_m ou r_M puisque $\vec{V}_0 \perp \vec{OM}_0$ et donc $\dot{r}_0 = 0$. Pour que ce point soit le plus éloigné, il faut que $D > r_1$ soit avec la relation entre r_1 et F , on trouve $F > mV_0^2/D$. Le point ne peut atteindre O car il faudrait une énergie mécanique infinie.

- d) Pour avoir un mouvement circulaire, il faut $D = r_1$ soit $F = mV_0^2/D$. Dans ce cas, le mouvement est circulaire et forcément uniforme car $\dot{\theta} = C/r_1^2 = cte$

A.4 On étudie la perle dans le référentiel terrestre considéré comme galiléen. Elle est soumise à son poids $m\vec{g}$ et à la réaction $\vec{R}$ du cerceau, cette réaction est dirigée selon $\vec{u}_r$. L'application du théorème du moment cinétique donne $\frac{d\vec{L}_O}{dt} = \vec{OM} \wedge m\vec{g} + \vec{OM} \wedge \vec{R}$. Le deuxième terme est nul car les deux vecteurs sont colinéaires. On en déduit en explicitant les expressions du poids et du moment cinétique $m r^2 \dot{\theta} \vec{u}_z = -mg \sin \theta \vec{u}_z$ soit $\ddot{\theta} + \frac{g}{r} \sin \theta = 0$.

A.5 1. On étudie le skieur dans le référentiel terrestre considéré comme galiléen. Il est soumis à son poids $m\vec{g}$ et à la réaction $\vec{R}$ de la piste qui est normale du fait de l'absence de frottement. L'application du théorème du moment cinétique par rapport à O le centre de la piste s'écrit $mR^2 \ddot{\alpha} \vec{u}_y = mgR \cos \alpha \vec{u}_y$ avec $\vec{u}_y$ un vecteur unitaire horizontal perpendiculaire au plan du mouvement soit $\ddot{\alpha} - \frac{g}{R} \cos \alpha = 0$.

2. En multipliant l'équation du mouvement par $\dot{\alpha}$ et en intégrant, on en déduit $\frac{\dot{\alpha}^2}{2} = \frac{g}{R} (\sin \alpha - \sin \alpha_0)$, ce qui permet d'obtenir l'expression de la vitesse

$$v = R\dot{\alpha} = \sqrt{2gR(\sin \alpha - \sin \alpha_0)}$$

3. On a une vitesse maximale lorsque $\sin \alpha$ est maximum soit $\alpha = 90^\circ$. La valeur de la vitesse maximale est donc $v_{\text{max}} = 29,5 \text{ m.s}^{-1} = 106,3 \text{ km.h}^{-1}$.

B.1 1. L'énergie mécanique est la somme de l'énergie potentielle et de l'énergie cinétique. Dans le cas d'un potentiel central, la force est colinéaire à $\vec{OM}$, donc le moment de cette force par rapport à O est nul. Par application du théorème du moment cinétique, on en déduit que le moment cinétique est un vecteur constant. La vitesse et le vecteur $\vec{OM}$ sont perpendiculaires à $\vec{\sigma}_0$, donc à un vecteur constant et donc toujours dans le même plan : le mouvement est plan et on utilise les coordonnées polaires.

La vitesse en coordonnées polaires est : $\vec{v} = \dot{r}\vec{u}_r + r\dot{\theta}\vec{u}_\theta$ donc l'expression générale du moment cinétique en coordonnées polaires est :

$$\vec{\sigma}_O = \vec{OM} \wedge m\vec{v} = mr\vec{u}_r \wedge (\dot{r}\vec{u}_r + r\dot{\theta}\vec{u}_\theta) = mr^2\dot{\theta}\vec{u}_z$$

On en déduit l'expression de E_m dans laquelle on remplace $\dot{\theta}$ par $\sigma_0/(mr^2)$:

$$E_m = \frac{1}{2}mv^2 + E_p = \frac{1}{2}m(\dot{r}^2 + r^2\dot{\theta}^2) - \frac{Km}{r} = \frac{1}{2}m\dot{r}^2 + \frac{\sigma_0^2}{2mr^2} - \frac{Km}{r}$$

L'énergie potentielle efficace est $E_{p,\text{eff}} = \frac{\sigma_0^2}{2mr^2} - \frac{Km}{r}$.

On a donc $E_m = \frac{1}{2}m\dot{r}^2 + E_{p,\text{eff}}$. Étant donné que $\dot{r}^2 \geq 0$, le mouvement ne peut donc avoir lieu que si $E_m \geq E_{p,\text{eff}}$.

2. On a donc :
$$E_m \geq \frac{\sigma_0^2}{2mr^2} - \frac{Km}{r} \Rightarrow 2mr^2 E_m + 2m^2 r K - \sigma_0^2 \geq 0$$

Le problème consiste donc à déterminer l'intervalle dans lequel le trinôme est positif. Le coefficient du terme de plus haut degré est négatif (cas $E_m \leq 0$) : le trinôme sera positif pour $r \in [r_1, r_2]$ où r_1 et r_2 sont les racines. Il faut cependant pour que le problème ait une solution qu'il existe deux racines réelles positives. Le discriminant est $\Delta = m^4 K^2 + 2E_m \sigma_0^2 m$; il est positif si $E_m \geq -\frac{m^3 K^2}{2\sigma_0^2}$. On pourrait vérifier que cette valeur est le minimum de $E_{p,\text{eff}}$. Le produit des racines égal à $-\sigma_0^2/(2mE_m)$ est positif et leur somme égale à $-mK/(E_m)$ aussi, donc les deux racines sont positives et égales à :

$$r_1 = -\frac{mK}{2E_m} - \frac{\sqrt{\Delta}}{4mE_m} \text{ et } r_2 = -\frac{mK}{2E_m} + \frac{\sqrt{\Delta}}{4mE_m}$$

3. L'inégalité précédente pour E_m (toujours négatif) entraîne :

$$E_m \geq -\frac{m^3 K^2}{2\sigma_0^2} \Rightarrow \sigma_0 \leq \sqrt{\frac{-m^3 K^2}{2E_m}} \text{ soit } \sigma_m = \sqrt{\frac{-m^3 K^2}{2E_m}}$$

Pour $\sigma_0 = \sigma_m$, le discriminant est nul donc le trinôme a une racine double r_d . Il se met donc sous la forme $2mE_m(r - r_d)^2$. Comme mE_m est négatif, la seule possibilité pour que le trinôme soit positif ou nul est qu'il soit nul. Ainsi $r = r_1$ et la trajectoire est un cercle.

4. Si $E_m = 0$, l'inégalité devient :

$$0 \geq \frac{\sigma_0^2}{2mr^2} - \frac{Km}{r} \Rightarrow r \geq \frac{\sigma_0^2}{2Km^2} \text{ soit } r_m = \frac{\sigma_0^2}{2Km^2}$$

en excluant le cas non physique $r = 0$

B.2 1. Le point se déplace dans le plan $(O, \vec{u}_r, \vec{u}_\theta)$ à distance L fixe de O donc la vitesse est $\vec{v} = L\dot{\theta}\vec{u}_\theta$. Par définition du moment cinétique, on peut écrire :

$$\vec{\sigma}_0 = \vec{OM} \wedge m\vec{v} \Rightarrow \vec{\sigma}_0 = L\vec{u}_r \wedge mL\dot{\theta}\vec{u}_\theta = mL^2\dot{\theta}\vec{u}_z$$

2. a) Les trois forces s'exerçant sur le point M sont : le poids, la tension du fil $\vec{T}$ et la réaction normale au plan incliné $\vec{R}_N$.

La force $\vec{T}$ est alignée avec le fil donc $\vec{T} = -T\vec{u}_r$ si l'on note T la norme de $\vec{T}$. La force $\vec{R}_N$ est perpendiculaire au plan donc selon $\vec{u}_z$ soit $\vec{R}_N = R_N\vec{u}_z$. Le poids est vertical donc dans le plan xOz soit $m\vec{g} = mg(\sin\alpha\vec{u}_x - \cos\alpha\vec{u}_z)$. Il faut ensuite projeter la composante dans le plan sur $\vec{u}_r$ et $\vec{u}_\theta$ sachant que $\vec{u}_x = \cos\theta\vec{u}_r - \sin\theta\vec{u}_\theta$. Finalement on trouve :

$$m\vec{g} = mg(\sin\alpha\cos\theta\vec{u}_r - \sin\alpha\sin\theta\vec{u}_\theta - \cos\alpha\vec{u}_z)$$

b) Le théorème du moment cinétique en O entraîne :

$$\frac{d\vec{\sigma}_0}{dt} = mL^2\ddot{\theta}\vec{u}_z = \vec{OM} \wedge m\vec{g} + \vec{OM} \wedge \vec{R}_N + \vec{OM} \wedge \vec{T}$$

Le moment de la force $\vec{T}$ par rapport à O est nul puisque $\vec{T}$ est colinéaire à $\vec{OM}$. Le moment de $\vec{R}_N$ est perpendiculaire à Oz . Il reste à calculer les composantes du moment du poids sur Oz :

$$\vec{OM} \wedge m\vec{g} = mgL(\cos\alpha\vec{u}_\theta - \sin\alpha\sin\theta)\vec{u}_z$$

Le théorème du moment cinétique s'écrit :

$$mL^2\ddot{\theta} = -mgL\sin\alpha\sin\theta \Rightarrow \ddot{\theta} + \frac{g\sin\alpha}{L}\sin\theta = 0$$

Pour les petites oscillations $\sin \theta \simeq \theta$ et si l'on pose $\omega_0^2 = \frac{g \sin \alpha}{L}$, on trouve l'équation de l'oscillateur harmonique : $\ddot{\theta} + \omega_0^2 \theta = 0$.

- c) La solution à l'équation précédente est $\theta = \theta_0 \cos(\omega_0 t + \varphi)$. À $t = 0$, $\theta = 0$ entraîne $\cos \varphi = 0$. On prend par exemple $\varphi = \pi/2$, alors $\theta = -\theta_0 \sin \omega_0 t$. Le calcul de $\dot{\theta}$ donne $\dot{\theta} = -\omega_0 \theta_0 \cos \omega_0 t$. À $t = 0$, $\dot{\theta} = v_0/L$ d'où $\theta_0 = -v_0/(L\omega_0)$ et

$$\theta(t) = \frac{v_0}{L\omega_0} \sin \omega_0 t$$

L'angle maximal atteint est donc $\frac{v_0}{L\omega_0}$.

- B.3** 1. Les trois forces qui s'exercent sur le point M sont le poids et la réaction normale $\vec{R}_n$ du plan xOy ainsi que la tension du ressort qui est selon OM . Puisque le mouvement est plan et horizontal, les forces verticales se compensent : $\vec{R}_n + m\vec{g} = \vec{0}$. Le théorème du moment cinétique s'écrit donc :

$$\frac{d\vec{L}_O}{dt} = \vec{OM} \wedge \vec{T} = \vec{0}$$

puisque $\vec{T}$ et $\vec{OM}$ sont colinéaires. Ainsi $\vec{L}_O$ est une constante.

2. a) L'expression générale du moment cinétique est $\vec{L}_O = \vec{OM} \wedge m\vec{v}$. À $t = 0$, la vitesse est nulle donc $\vec{L}_O = \vec{0}$. Ainsi $\vec{L}_O$ est nul à tout instant, donc la vitesse est toujours colinéaire à $\vec{OM}$ et le mouvement est rectiligne suivant l'axe Ox .
b) La projection de la relation fondamentale de la dynamique sur Ox donne $m\ddot{x} = -k(\ell - \ell_0)$. Or $x = \ell$ donc on obtient : $\ddot{x} + \omega_0^2 x = \omega_0^2 \ell_0$ avec $\omega_0^2 = k/m$. Les solutions sont sinusoïdales :

$$x(t) = A \cos \omega_0 t + B \sin \omega_0 t + \ell_0 \text{ et } \dot{x}(t) = -A\omega_0 \sin \omega_0 t + B\omega_0 \cos \omega_0 t$$

Les conditions initiales $x(0) = 1, 2\ell_0$ et $\dot{x}(0) = 0$ entraînent $A = 0, 2\ell_0$ et $B = 0$. L'intervalle de variation de ℓ est donc $[0, 8\ell_0, 1, 2\ell_0]$.

3. a) L'expression générale du moment cinétique en coordonnées polaires est :

$$\vec{L}_O = \vec{OM} \wedge m\vec{v} = mr\vec{u}_r \wedge (\dot{r}\vec{u}_r + r\dot{\theta}\vec{u}_\theta) = mr^2\dot{\theta}\vec{u}_z$$

À $t = 0$, $\vec{L}_O = \ell_1 \vec{u}_x \wedge m\ell_1 \omega \vec{u}_y$ soit $\vec{L}_O = m\ell_1^2 \omega \vec{u}_z$.

- b) L'énergie potentielle élastique est $E_p = \frac{k}{2}(\ell - \ell_0)^2$.

Dans le cas présent, le poids ne travaille pas donc l'énergie potentielle de poids est constante ; on n'a pas besoin d'en tenir compte. La réaction perpendiculaire au mouvement comme le poids ne travaille pas non plus donc le théorème de l'énergie cinétique donne :

$$dE_c = \delta W = -dE_p \Rightarrow d(E_c + E_p) = 0 \Rightarrow E_m = cte$$

L'expression de E_m en fonction des conditions initiales est :

$$E_m = \frac{m}{2}\ell_1^2 \omega^2 + \frac{k}{2}(\ell_1 - \ell_0)^2$$

L'expression générale de E_m est :

$$E_m = \frac{1}{2}mv^2 + E_p = \frac{1}{2}(\dot{r}^2 + r^2\dot{\theta}^2) + \frac{1}{2}k(r - \ell_0)^2$$

- c) Puisque le moment cinétique est constant et égal à $mr^2\dot{\theta}$, la vitesse angulaire est de signe constant ici positif en raison des conditions initiales. Donc le module L est égal à $mr^2\dot{\theta}$ et $\dot{\theta} = L/mr^2$. On remplace $\dot{\theta}$ par cette expression dans l'énergie mécanique et on trouve :

$$E_m = \frac{1}{2}m\dot{r}^2 + \frac{L^2}{2mr^2} + \frac{1}{2}k(r - \ell_0)^2 \text{ d'ou } E_{\text{eff}}(r) = \frac{L^2}{2mr^2} + \frac{1}{2}k(r - \ell_0)^2$$

On a représenté l'allure de la fonction E_{eff} sur la figure (S25.3).

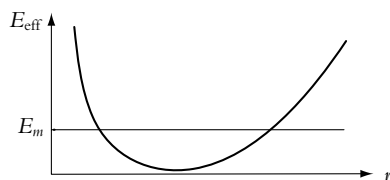


Figure S25.3

Le mouvement ne peut avoir lieu que dans les zones où $E_m \geq E_{\text{eff}}$ puisque $m\dot{r}^2 \geq 0$, donc la zone de mouvement est sous la droite d'équation $E_{\text{eff}} = E_m$.

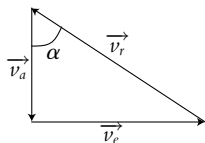
- d) D'après ce que l'on vient d'écrire, il faudrait une énergie mécanique infinie pour s'éloigner à l'infini de O .
- e) La vitesse radiale $\dot{r}$ peut s'annuler au cours du mouvement aux deux points extrêmes pour lesquels $E_m = E_{\text{eff}}$. En revanche la vitesse angulaire ne peut s'annuler car avec la constance de L cela entraînerait $r \rightarrow \infty$ ce qui est impossible. La vitesse ne peut donc pas s'annuler.
- f) Il faudrait aussi une énergie infinie pour atteindre O .
4. a) Si le mouvement est circulaire r est constant donc $\dot{\theta} = L/mr^2$ est aussi constante.
- b) La seule possibilité correspond au minimum de la courbe $E_{\text{eff}}(r)$ sinon comme c'est représenté sur la figure (S25.3), le point M oscille entre deux positions. Il faut donc calculer la position du minimum :

$$\frac{dE_{\text{eff}}}{dr}(\ell_1) = -\frac{L^2}{m\ell_1^3} + k(\ell_1 - \ell_0) = 0 \Rightarrow -\frac{m^2\ell_1^4\omega^2}{m\ell_1^3} + k(\ell_1 - \ell_0) = 0$$

On trouve $\ell_1 = \ell_0 \frac{\omega^2}{\omega_0^2 - \omega^2}$, avec $\omega_0 = \sqrt{k/m}$. Cette solution est possible si $\omega < \omega_0$.

Chapitre 26

A.1 La loi de composition des vitesses s'écrit $\vec{v}_a = \vec{v}_e + \vec{v}_r$ en notant $\vec{v}_a$ la vitesse de la neige par rapport au sol, $\vec{v}_e$ la vitesse de la voiture par rapport au sol et $\vec{v}_r$ la vitesse de la neige par rapport à la voiture. Graphiquement cela se traduit par



D'après les données de l'énoncé, on a $\alpha = 80^\circ$, $v_e = 110 \text{ km.h}^{-1}$. On souhaite déterminer les valeurs de v_a et v_r .

Le schéma précédent permet d'écrire $\tan \alpha = \frac{v_e}{v_a}$ soit $v_a = \frac{v_e}{\tan \alpha} = 19,4 \text{ km.h}^{-1}$.

Par ailleurs, en appliquant la relation de Pythagore, on a $v_r = \sqrt{v_a^2 + v_e^2} = 112 \text{ km.h}^{-1}$.

A.2 On note $\mathcal{R}$ le référentiel absolu lié au sol et $\mathcal{R}_t$ celui lié au tapis. Le référentiel $\mathcal{R}_t$ est en translation par rapport à $\mathcal{R}$. La loi de composition des vitesses donne :

$$\vec{V}_a = \vec{V}_{/\mathcal{R}} = \vec{V}_{/\mathcal{R}_t} + \vec{V}_e \text{ soit } \vec{V}_a = \vec{V} + \vec{V}_t$$

1. On a représenté sur la figure (S26.1) les trois vecteurs vitesses. Il faut calculer l'angle du vecteur $\vec{V}_a$ par rapport aux bords du tapis, soit $\tan \alpha = V_t/V$. Sur la figure, le point de départ est A et le point d'arrivée B . La distance d parcourue entre les deux est $d = a \tan \alpha$ soit $d = aV_t/V$.

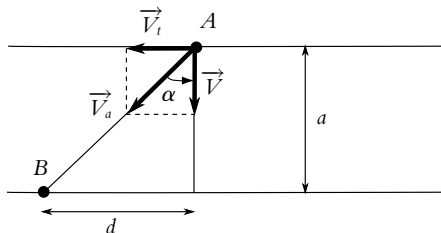


Figure S26.1

Le temps de traversée est $t_1 = AB/V_a$, or $AB = a/\cos \alpha$ et $V_a = V/\cos \alpha$ donc $t_1 = a/V$.

2. La vitesse absolue $\vec{V}_a$ doit maintenant être perpendiculaire aux bords du tapis (figure S26.2).

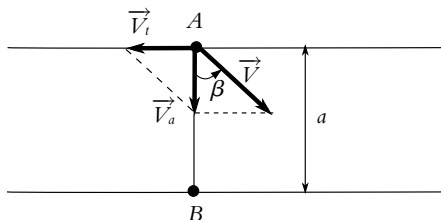


Figure S26.2

L'angle β que doit faire $\vec{V}$ avec la normale aux bords du tapis est égal à $\text{Arcsin}\left(\frac{V_t}{V}\right)$. Le temps de traversée est $t_2 = a/V_a$ soit $t_1 = a/\sqrt{V^2 - V_t^2}$. Cette trajectoire est possible si $V > V_t$.

3. Dans le repère xOy de la figure (S26.3) la vitesse $\vec{V}_a$ s'écrit :

$$\vec{V}_a = (V_t + V \cos \theta) \vec{u}_x + V \sin \theta \vec{u}_y$$

Le temps de traversée est donc $t_3 = a/V_{a,y}$ soit $t_3 = a/(V \sin \theta)$. Le temps t_3 est donc minimal pour $\theta = \pi/2$ ce qui revient au cas de la première question.

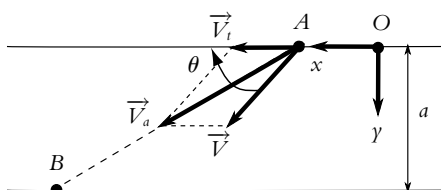
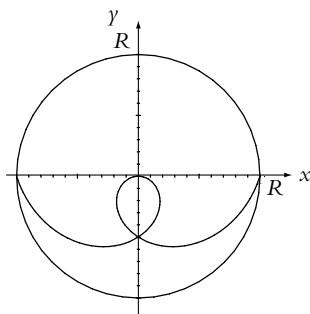


Figure S26.3

- A.3** 1. On repère la position x de l'enfant sur la droite AB orientée de A vers B en prenant O comme origine. La vitesse vaut $v = \frac{dx}{dt}$ et initialement $x(0) = -R$ donc $x = -R + vt$. La trajectoire est la droite AB .

2. Les coordonnées polaires sont $x = -R + vt$ et $\theta = \omega t$. On en déduit $t = \frac{\theta}{\omega}$ et $x = -R + \frac{v}{\omega} \theta$.

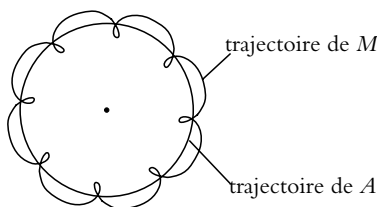
3. Au bout d'un tour, on a $x = R$ pour $\theta = 2\pi$ soit en reportant dans l'équation de la trajectoire dans le référentiel lié au sol $v = \frac{R\omega}{\pi}$ et $x = R\left(\frac{\theta}{\pi} - 1\right)$. Le calcul de la dérivée de x par rapport à θ donne $x'(\theta) = \frac{R}{\pi}$ donc x est une fonction croissante de θ avec $x(0) = -R$, $x(\pi) = 0$ et $x(2\pi) = R$. On en déduit l'allure suivante :



4. On obtient la vitesse et l'accélération dans le référentiel lié au sol soit en dérivant $\overrightarrow{OM} = x\vec{u}_r$ par rapport au temps soit en utilisant les lois de compositions des vitesses et des accélérations à partir de leur expression dans le référentiel lié au manège. Dans cette deuxième hypothèse, on a une vitesse relative $\vec{v}_r = v\vec{u}_r$, une vitesse d'entraînement $\vec{v}_e = x\dot{\theta}\vec{u}_\theta$, une accélération relative nulle, une accélération d'entraînement $\vec{a}_e = -\omega^2 x\vec{u}_r$ et une accélération de Coriolis $\vec{a}_c = 2v\dot{\theta}\vec{u}_\theta$. On en déduit $\vec{v} = v\vec{u}_r + x\dot{\theta}\vec{u}_\theta$ et $\vec{a} = 2v\dot{\theta}\vec{u}_\theta - x\dot{\theta}^2\vec{u}_r$.

A.4 A est en rotation autour de O dans le référentiel fixe : il décrit un cercle de centre O et de rayon OA .

M décrit un cercle de centre A et de rayon AM dans le référentiel lié à la barre. Dans le référentiel fixe, M décrit une spirale circulaire :



A.5 La loi de composition des vitesses s'écrit $\vec{v}_a = \vec{v}_r + \vec{v}_e$ avec $\vec{v}_a$ la vitesse dans le référentiel fixe, $\vec{v}_r$ la vitesse dans le référentiel mobile et la vitesse du référentiel mobile par rapport au référentiel fixe $\vec{v}_e = v_0\vec{u}_x$. Pour obtenir la vitesse $\vec{v}_r$, on peut utiliser les coordonnées du point M dans le référentiel mobile soit $\overrightarrow{OM} = R\cos\theta\vec{u}_x + R\sin\theta\vec{u}_y$ avec $\theta = \omega_0 t$ puis dériver par rapport au temps soit $\vec{v}_r = -R\omega_0\sin\omega_0 t\vec{u}_x + R\omega_0\cos\omega_0 t\vec{u}_y$. On en déduit la vitesse dans le référentiel fixe $\vec{v}_a = (v_0 - R\omega_0\sin\omega_0 t)\vec{u}_x + R\omega_0\cos\omega_0 t\vec{u}_y$.

En intégrant par rapport au temps, on obtient les coordonnées de M dans le référentiel fixe soit $\overrightarrow{OM} = (v_0 t + R\cos\omega_0 t)\vec{u}_x + R\sin\omega_0 t\vec{u}_y$.

La trajectoire est donc la courbe paramétrée définie par

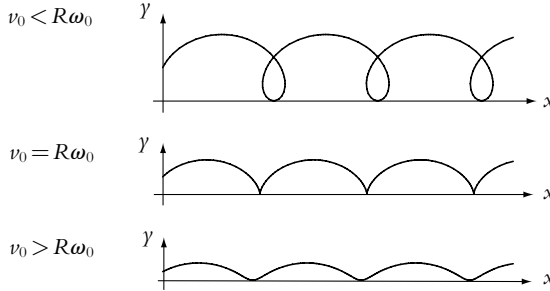
$$\begin{cases} x = v_0 t + R\cos\omega_0 t \\ y = R\sin\omega_0 t \end{cases}$$

L'étude de la fonction $x(t)$ montre que la dérivée $\dot{x} = v_0 - R\omega_0\sin\omega_0 t$ ne peut s'annuler si $v_0 > R\omega_0$ puisque le sinus est compris entre -1 et 1 . Dans ce cas, la dérivée est toujours positive donc x est

croissante. Dans le cas contraire, la dérivée peut s'annuler et changer de signe périodiquement donnant lieu à des extrema pour $x(t)$.

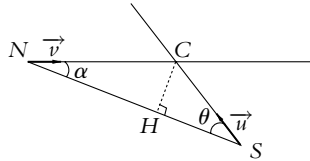
Quant à la fonction $y(t)$, sa dérivée s'annule périodiquement donnant lieu à des extrema de $y(t)$

On en déduit les allures suivantes :



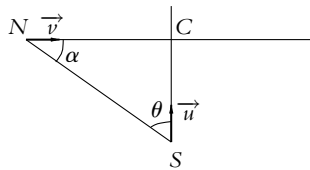
On note que dans le cas $v_0 = R\omega_0$, on obtient une cycloïde.

B.1 1. Les vitesses sont supposées uniformes donc : $NC = vt$ et $SC = ut$ en notant C le point d'impact s'il existe.



Alors d'après le schéma ci-dessus, $CH = NC \sin \alpha = SC \sin \theta$ soit un angle de tir : $\theta = \text{Arcsin} \left(\frac{v}{u} \sin \alpha \right)$.

2. On cherche à obtenir une distance SC minimale, ce qui sera le cas si le triangle SCN est rectangle en C et $\sin \theta_{\min} = \frac{NC}{SN} = \frac{v}{\sqrt{u^2 + v^2}}$.



B.2 1. On applique la loi de composition des vitesses en notant B le bac, $\mathcal{R}$ le référentiel lié au rivage et $\mathcal{R}'$ celui lié au fleuve et en translation par rapport à $\mathcal{R}$. On a : $\vec{v}_{/\mathcal{R}}(B) = \vec{v}_{/\mathcal{R}'}(B) + \vec{v}_e = \vec{v} + \vec{u}$. Dans le premier cas (le bac traverse perpendiculairement), $\vec{v}_{/\mathcal{R}}(B)$ et $\vec{u}$ sont perpendiculaires donc $v_{/\mathcal{R}}(B) = \sqrt{v^2 - u^2}$ et $T_1 = \frac{d}{\sqrt{v^2 - u^2}}$.

Dans le deuxième cas (le bac se déplace dans le sens du courant) : $T_2 = \frac{d}{u + v}$.

Dans le dernier cas (le bac se déplace dans le sens opposé au courant) : $T_3 = \frac{d}{v - u}$. Implicitement cela suppose que $v > u$.

On en déduit :

$$\frac{T_1}{T_2} = \frac{d(u + v)}{d\sqrt{v^2 - u^2}} = \sqrt{\frac{v + u}{v - u}} \quad \text{et} \quad \frac{T_3}{T_1} = \frac{d\sqrt{v^2 - u^2}}{d(v - u)} = \sqrt{\frac{v + u}{v - u}} = \frac{T_1}{T_2}$$

$$2. \frac{T_1}{T_2} = \sqrt{\frac{v+u}{v-u}} = 3 \text{ soit } v = \frac{5}{4}u.$$

$$3. \text{ Dans le premier cas : } \alpha = \text{Arcsin} \frac{u}{v} = \text{Arcsin} \frac{4}{5} = 53^\circ \text{ et } v_{/\mathcal{R}}(B) = \sqrt{v^2 - u^2} = \frac{3}{4}u$$

$$\text{Dans le second : } v(B)_{/\mathcal{R}} = v + u = \frac{9}{4}u.$$

$$\text{Dans le dernier cas : } v(B)_{/\mathcal{R}} = v - u = \frac{1}{4}u$$

$$4. \frac{T_3}{T_2} = \frac{T_3}{T_1} \frac{T_1}{T_2} = \left(\frac{T_1}{T_2} \right)^2 \text{ donc } T_3 = 9T_2 \text{ qu'on peut aussi obtenir à partir du rapport des vitesses entre les deuxième et troisième cas.}$$

B.3 1. Par rapport à $\mathcal{R}_S$ le référentiel $\mathcal{R}_C$ à un mouvement de rotation de vecteur rotation $\vec{\Omega} = \dot{\theta} \vec{u}_z$ et d'axe Oz si l'on note $\vec{u}_z$ le troisième vecteur de la base cylindrique $(\vec{u}_r, \vec{u}_\theta, \vec{u}_z)$.

2. Dans $\mathcal{R}_C$ le point M a un mouvement circulaire de vitesse angulaire $\dot{\varphi}$ donc si l'on définit une base polaire $(\vec{u}_\rho, \vec{u}_\varphi)$ de centre C (figure S26.4), on peut écrire :

$$\vec{v}_{\mathcal{R}_C} = b \dot{\varphi} \vec{u}_\varphi = b \dot{\varphi} (\cos \varphi \vec{u}_r + \sin \varphi \vec{u}_\theta)$$

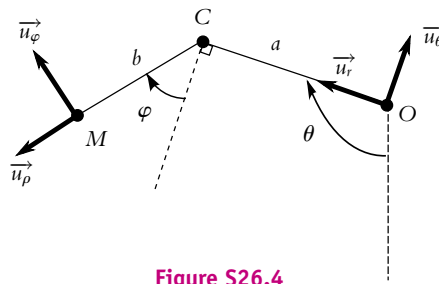


Figure S26.4

3. On utilise la relation de Chasles $\vec{OM} = \vec{OC} + \vec{CM}$. On a $\vec{OM} = a \vec{u}_r$ et $\vec{CM} = b \vec{u}_\rho$. De plus, dans la base $(\vec{u}_r, \vec{u}_\theta)$, on peut écrire :

$$\vec{u}_\rho = \sin \varphi \vec{u}_r - \cos \varphi \vec{u}_\theta \text{ soit } \vec{OM} = (a + b \sin \varphi) \vec{u}_r - b \cos \varphi \vec{u}_\theta$$

La vitesse d'entraînement est la vitesse, par rapport à $\mathcal{R}_S$, du point fixe P de $\mathcal{R}_C$ coïncidant avec M . Ce point, a par rapport à O , un mouvement circulaire de rayon OP et de vitesse angulaire $\dot{\theta}$. Cela revient à fixer φ puisque P est fixe par rapport à C . On dérive donc l'expression de $\vec{OM}$ par rapport à t en gardant φ constant :

$$\vec{v}_e = \frac{d\vec{OM}}{dt} = (a + b \sin \varphi) \frac{d\vec{u}_r}{dt} - b \cos \varphi \frac{d\vec{u}_\theta}{dt} = \dot{\theta} b \cos \varphi \vec{u}_r + \dot{\theta} (a + b \sin \varphi) \vec{u}_\theta$$

On peut aussi calculer la vitesse avec l'expression : $\vec{v}_e = \vec{\Omega} \wedge \vec{OM}$.

4. Pour calculer la vitesse $\vec{v}_{/\mathcal{R}_S}$ dans $\mathcal{R}_S$, on applique la loi de composition des vitesses soit :

$$\vec{v}_{/\mathcal{R}_S} = \vec{v}_{/\mathcal{R}_C} + \vec{v}_e = b \cos \varphi (\dot{\varphi} + \dot{\theta}) \vec{u}_r + [\dot{\theta} (a + b \sin \varphi) + b \dot{\varphi} \sin \varphi] \vec{u}_\theta$$

Lors du lâcher ($\theta = \pi$ et $\varphi = \pi/2$), l'expression précédente devient :

$$\vec{v}_{/\mathcal{R}_S} = +[\dot{\theta}(a+b) + b\dot{\varphi}] \vec{u}_\theta$$

S'il n'y avait qu'un seul bras rigide, le point M aurait un mouvement circulaire de rayon $a+b$ et de vitesse angulaire $\dot{\theta}$. Sa vitesse serait $(a+b)\dot{\theta}$. Avec l'articulation, on a donc un terme supplémentaire $b\dot{\varphi}$ avec $\dot{\varphi} > 0$. Un système du même genre était aussi utilisé par les hommes préhistoriques pour

lancer des javelots avec le bras. Il repliait une partie articulée qui se déplaçait au moment du lancer avec le mouvement du bras.

B.4 1. a) La vitesse $\vec{V}$ est la vitesse de l'avion dans le référentiel de l'air et $\vec{v}_e$ est la vitesse d'entraînement. Par la composition des vitesses, on trouve $\vec{v}_a = \vec{V} + \vec{v}_e$.

b) Pour qu'il se déplace en ligne droite, il faut que la projection de $\vec{v}_a$ perpendiculairement à AB soit nulle, soit :

$$v_e \sin \varphi - V \sin \theta = 0$$

c) On a alors $\sin \theta = v_e \frac{\sin \varphi}{V}$ soit $\sin \theta = 0,043$ et $\theta = 2,47^\circ$.

2. a) On a représenté les vitesses à l'aller et au retour sur la figure (S26.5).

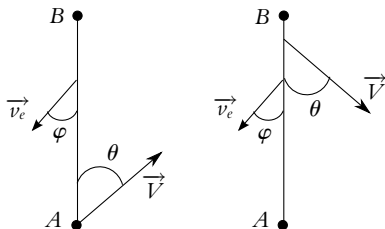


Figure S26.5

À l'aller, la projection de la vitesse $\vec{v}_a$ sur AB (dans le sens de AB) donne $V \cos \theta - v_e \cos \varphi$ donc pour l'aller le temps mis est $t_a = d / (V \cos \theta - v_e \cos \varphi)$.

Pour le retour, la projection de la vitesse dans le sens de BA est $V \cos \theta + v_e \cos \varphi$. Le temps mis pour le retour $t_r = d / (V \cos \theta + v_e \cos \varphi)$. Le temps mis pour l'aller-retour est :

$$t_{ar} = d \left(\frac{1}{V \cos \theta - v_e \cos \varphi} + \frac{1}{V \cos \theta + v_e \cos \varphi} \right) = 2,28 \text{ h}$$

b) L'aller retour sans vent correspondrait à $v_e = 0$ et $\theta = 0$. Le temps mis serait $t'_{ar} = 2d/V = 2,24$ h soit 2,4 minutes de différence. Si l'on compare le temps de l'aller avec et sans vent on trouve 1,28 h et 1,12 h donc l'avion met plus de temps car il est contre le vent. Pour le retour le temps sans vent est le même qu'à l'aller (1,12 h) et le temps avec vent est 1 h donc on ne rattrape pas au retour tout le temps perdu à l'aller.

On voit sur l'expression que les termes dus au vent ne peuvent se compenser car ils sont au dénominateur. On peut exprimer la relation entre t_{ar} et t'_{ar} en réduisant au même dénominateur l'expression de t_{ar} :

$$t_{ar} = d \frac{2V \cos \theta}{V^2 \cos^2 \theta - v_e^2 \cos^2 \varphi} = \frac{2d}{V} \frac{\cos \theta}{\cos^2 \theta - \left(\frac{v_e}{V}\right)^2 \cos^2 \varphi} = t'_{ar} \frac{\cos \theta}{\cos^2 \theta - \left(\frac{v_e}{V}\right)^2 \cos^2 \varphi}$$

Pour v_e/V fixé on va étudier comment varie le rapport, pour cela on l'exprime en fonction de φ en utilisant $\sin \theta = v_e \frac{\sin \varphi}{V}$:

$$\frac{\cos \theta}{\cos^2 \theta - \left(\frac{v_e}{V}\right)^2 \cos^2 \varphi} = \frac{\sqrt{1 - \sin^2 \theta}}{1 - \sin^2 \theta - \left(\frac{v_e}{V}\right)^2 \cos^2 \varphi} = \frac{\sqrt{1 - \left(\frac{v_e}{V}\right)^2 \sin^2 \varphi}}{1 - \left(\frac{v_e}{V}\right)^2} = f(\varphi)$$

L'angle φ varie entre 0 et $\pi/2$, intervalle sur lequel $\sin \varphi$ est croissante et $f(\varphi)$ décroissante. Pour $\varphi = \pi/2$, la fonction $f(\varphi) \geq 1$ donc sur l'intervalle elle est toujours supérieure ou égale à 1 et $t_{ar} \geq t'_{ar}$. Dans le cas où le vent est complètement latéral, les deux temps sont égaux.

3. Il faut calculer les temps t_{CD} et t_{EB} en utilisant les résultats de la question précédente. Le temps t_{CD} correspond au temps du retour (sens du vent) et t_{EB} au temps de l'aller (contre le vent). On ne connaît pas la distance $CD = EB$ que l'on note d' .

On doit calculer le nouvel angle $\theta = \text{Arcsin}(v_e \frac{\sin \varphi}{V})$ soit $\theta = 4,94^\circ$. On peut déterminer le rapport $r = t_{CD}/t_{EB}$:

$$r = \frac{t_{CD}}{t_{EB}} = \frac{V \cos \theta - v_e \cos \varphi}{V \cos \theta + v_e \cos \varphi} = 0,39$$

On connaît d'autre part $t_{CD} + t_{EB} = 120$ s car chaque virage prend 1 minute et l'ensemble du trajet prend 4 minutes. On en déduit $t_{EB} = 86,3$ s et $t_{CD} = 33,7$ s. Le virage DE doit être amorcé 93,7 s après le premier passage en B .

On peut calculer la distance $d' = (V \cos \theta + v_e \cos \varphi)t_{EB}$, soit $d' = 6,65$ km.

- B.5** 1. Le point C a un mouvement circulaire uniforme de centre O et de vitesse $\vec{v}_C = v_0 \vec{u}_\varphi$. La vitesse angulaire est égale à :

$$\omega_1 = \frac{d\varphi}{dt} = \frac{v_0}{OC} = \frac{v_0}{R-r} \Rightarrow \varphi = \frac{v_0}{R-r}t + cte$$

À $t = 0$, $\varphi = 0$ donc la constante est nulle.

2. Le point M a un mouvement circulaire de centre C et de vitesse angulaire ω_2 dans $\mathcal{R}'$. Sa vitesse $\vec{v}_{/R'}$ est $\omega_2 \vec{u}_z \wedge \vec{CM}$ où $\vec{u}_z$ est le troisième vecteur de la base polaire. On obtient après projection sur $(\vec{u}_r, \vec{u}_\varphi)$:

$$\vec{v}_{/R'} = \omega_2 \vec{u}_z \wedge r(\cos \theta \vec{u}_r + \sin \theta \vec{u}_\varphi) = r\omega_2(-\sin \theta \vec{u}_r + \cos \theta \vec{u}_\varphi)$$

3. On applique la relation de Chasles : $\vec{OM} = \vec{OC} + \vec{CM}$ ce qui donne :

$$\vec{OM} = ((R-r) + r \cos \theta) \vec{u}_r + r \sin \theta \vec{u}_\varphi$$

Le point P (fixe dans $\mathcal{R}'$) coïncidant avec le point M a un mouvement circulaire uniforme de vitesse angulaire ω_1 , de centre O dans $\mathcal{R}$, donc sa vitesse dans $\mathcal{R}$ qui est aussi la vitesse d'entraînement est :

$$\vec{v}_e = \omega_1 \vec{u}_z \wedge \vec{OM} = -\omega_1 r \sin \theta \vec{u}_r + \omega_1 ((R-r) + r \cos \theta) \vec{u}_\varphi$$

On applique la loi de composition des vitesses pour calculer la vitesse absolue (dans $\mathcal{R}$) :

$$\vec{v}_a = \vec{v}_e + \vec{v}_{/R'} = -r \sin \theta (\omega_1 + \omega_2) \vec{u}_r + (\omega_1 ((R-r) + r \cos \theta) + \omega_2 r \cos \theta) \vec{u}_\varphi$$

4. Le point I correspond à $\theta = 0$, donc sa vitesse dans $\mathcal{R}$ qui doit être nulle est :

$$\vec{v}_a(I) = (\omega_1 ((R-r) + r) + \omega_2 r) \vec{u}_\varphi = \vec{0} \Rightarrow \omega_2 = -\omega_1 \frac{R}{r}$$

5. L'accélération d'entraînement est égale à l'accélération du point P coïncidant avec M , dans $\mathcal{R}$. Le point P , fixe dans $\mathcal{R}'$, a un mouvement circulaire uniforme dans $\mathcal{R}$ puisque d'après la relation précédente $\omega_1 = cte$. L'accélération dans ce cas est centripète et égale à $\vec{a}_e = -\omega_1^2 \vec{OM}$, soit :

$$\vec{a}_e = -\omega_1^2 ((R-r) + r \cos \theta) \vec{u}_r - r \sin \theta \vec{u}_\varphi$$

L'accélération de coriolis est obtenue par :

$$\vec{a}_c = 2\omega_1 \wedge \vec{v}_{/R'} = -2\omega_1 \omega_2 r (\cos \theta \vec{u}_r + \sin \theta \vec{u}_\varphi)$$

Enfin, l'accélération relative correspond au mouvement circulaire uniforme de M dans $\mathcal{R}'$ soit :

$$\vec{a}_r = -\omega_2^2 \vec{CM} = -\omega_2^2 r (\cos \theta \vec{u}_r + \sin \theta \vec{u}_\varphi)$$

L'accélération dans $\mathcal{R}$ est obtenue en sommant les trois vecteurs précédents (loi de composition des accélérations)

Chapitre 27

A.1 1. On étudie le singe dans le référentiel de la balançoire qui n'est pas galiléen. Il est soumis à son poids $m\vec{g}$, à la réaction de la balançoire et à la force d'inertie d'entraînement $-m\vec{a}_e$ (il n'y a pas de force d'inertie de Coriolis puisque le mouvement de la balançoire est un mouvement de translation). La projection du principe fondamental de la dynamique sur la verticale donne $m\ddot{z} = -mg + R - ma_z$ avec $\ddot{z} = 0$ puisque le singe est immobile par rapport à la balançoire. On en déduit $R = m(g + a_z)$.

2. Pour la phase d'accélération constante, le poids est surestimé puisque l'accélération d'entraînement est dans le même sens que l'accélération de pesantier.

Pour la phase à vitesse constante, le poids est bien estimé puisque l'accélération d'entraînement est nulle.

Pour la phase de décélération, le poids est sousestimé puisque l'accélération d'entraînement est de sens opposé à l'accélération de pesantier.

A.2 1. Le référentiel $\mathcal{R}_F$ lié à la fusée est en translation par rapport au référentiel géocentrique avec une accélération $\vec{a}_e$. La relation fondamentale de la dynamique appliquée à une personne P dans la fusée donne :

$$m\vec{a}/\mathcal{R}_F = \vec{f}_{ie} + \vec{f}_{ic} + \vec{A}_{T \rightarrow P} + \vec{R}$$

où la force d'inertie d'entraînement est $\vec{f}_{ie} = -m\vec{a}_e$, la force d'inertie de coriolis $\vec{f}_{ic}$ est nulle puisque $\mathcal{R}_F$ est en translation par rapport à $\mathcal{R}_G$, $\vec{A}_{T \rightarrow P}$ est l'attraction terrestre et $\vec{R}$ la réaction du sol de la fusée.

Dans la fusée les passagers sont immobiles donc $\vec{a}/\mathcal{R}_F = \vec{0}$.

Les passagers ressentent la même attraction qu'à la surface de la Terre ; cette force est l'opposée de la réaction du sol soit $-\vec{R} = -\frac{GmM_T}{R_T^2}\vec{u}_r$, où $\vec{u}_r$ est un vecteur unitaire orienté du centre de la Terre vers la fusée.

Si l'on projette la relation fondamentale sur $\vec{u}_r$ avec $a_e = \|\vec{a}_e\|$ et h l'altitude de la fusée par rapport au sol, on obtient :

$$0 = -ma_e - \frac{GmM_T}{(R_T + h)^2} + \frac{GmM_T}{R_T^2} \Rightarrow a_e = GM_T \left(\frac{1}{R_T^2} - \frac{1}{(R_T + h)^2} \right)$$

On remarque que la relation est valable quelle que soit la masse de la personne. On trouve qu'au départ, l'accélération doit être nulle ce qui est normale ; la fusée ne pourrait bien sûr pas décoller. Pour le décollage, en fait l'accélération est tellement importante que les personnages perdent connaissance. Pendant le trajet, la relation précédente nous indique que a_e dépende de la distance à la surface de la Terre ; il faut donc sans cesse ajuster la valeur de a_e ce qui semble difficile.

2. On choisit maintenant un vecteur $\vec{u}_r$ orienté du centre de la Lune vers la fusée. Comme il y a décélération $\vec{a}_e = a_e\vec{u}_r$. On veut toujours avoir $-\vec{R} = -\frac{GmM_T}{R_T^2}\vec{u}_r$ mais il faut prendre en compte l'attraction de la Lune sur P et l'on néglige l'attraction de la Terre sur P . On obtient donc :

$$0 = -ma_e - \frac{GmM_L}{(R_L + h)^2} + \frac{GmM_T}{R_T^2} \Rightarrow a_e = GM_T \left(\frac{1}{R_T^2} - \frac{1}{(R_L + h)^2} \right)$$

où h est l'altitude par rapport à la surface lunaire.

A.3 On note $\mathcal{R}$ le référentiel galiléen lié au sol et $\mathcal{R}_C$ le référentiel lié au camping-car.

1. La bille a donc une accélération dans le sens opposé au déplacement du camping-car. Le référentiel $\mathcal{R}_C$ est en translation par rapport au référentiel $\mathcal{R}$. La relation fondamentale de la dynamique dans $\mathcal{R}_C$ s'écrit donc :

$$m\vec{a}/\mathcal{R}_C = \vec{f}_{ie} + \vec{R} + m\vec{g}$$

$\vec{f}_{ie}$ étant la force d'inertie d'entraînement, $\vec{R}$ étant la réaction du support et sachant qu'on a négligé les frottements. Les deux forces verticales $\vec{R}$ et $m\vec{g}$ se compensent donc si $\vec{a}_{/R_C}$ est vers l'arrière c'est que $\vec{f}_{ie}$ aussi. L'expression de $\vec{f}_{ie}$ est $-m\vec{a}_e$ donc le véhicule a une accélération par rapport à $\mathcal{R}$ dans le sens de son déplacement.

2. On écrit de nouveau la relation fondamentale de la dynamique dans $\mathcal{R}_C$ avec $\vec{f}_{ie} = -m\vec{a}$:

$$m\vec{a}_{/R_C} = -m\vec{a} + \vec{R} + m\vec{g}$$

On projette l'équation sur un axe horizontal Ox orienté dans le sens de $\vec{a}$ (vers l'avant du véhicule) avec O position initiale de la bille ; le point O est donc lié au véhicule. On obtient : $m\ddot{x} = -ma$ soit en intégrant puisque la vitesse initiale et la position initiale dans $\mathcal{R}_C$ sont nulles $x = -at^2/2$. Le mouvement de la bille est rectiligne uniformément accéléré par rapport au camping-car.

- A.4** 1. Sur la durée du mouvement (quelques minutes), on peut considérer le référentiel terrestre comme un bon référentiel galiléen. Donc $\mathcal{R}$ est le référentiel terrestre. Le référentiel $\mathcal{R}'$ lié au manège est en mouvement de rotation d'axe vertical Oz et de vecteur rotation $\vec{\omega} = \omega\vec{u}_z$, le vecteur $\vec{u}_z$ étant le vecteur unitaire directeur de l'axe Oz .
2. La vitesse initiale (dans $\mathcal{R}'$) est dirigée vers la cible. La force d'inertie de Coriolis est égale à $-2m\vec{\omega} \wedge \vec{v}_{/R'}$ ce qui donne dès le départ une force dirigée vers la droite par rapport au lanceur. La balle frappera la paroi à droite de la cible.
3. La force de Coriolis est proportionnelle à la vitesse relative, donc si celle-ci est plus importante, la force a une intensité plus forte. Cependant, le temps de parcours sera plus petit donc l'effet de la déviation moins important. Il est difficile de conclure sans un calcul.
4. On néglige l'effet du poids sur la durée du parcours.
5. On définit un repère lié au manège $(0, \vec{u}_x, \vec{u}_y, \vec{u}_z)$. Le lanceur se trouve en O et la cible à la distance d de O sur l'axe Ox . Le système est la balle M dont on étudie le mouvement dans $\mathcal{R}'$. Les forces qui s'exercent sont uniquement les forces d'inertie puisqu'on néglige le poids. On suppose d'après l'énoncé que le mouvement est horizontal. La relation fondamentale dans $\mathcal{R}'$ s'écrit :

$$m\vec{a}_{/R'} = \vec{f}_{ie} + \vec{f}_{ic}$$

La force d'inertie d'entraînement est égale à $-m\vec{a}_e$, où $\vec{a}_e$ est l'accélération du point P (fixe dans $\mathcal{R}'$) coïncidant à M , par rapport à $\mathcal{R}$. Le mouvement de P dans $\mathcal{R}$ est circulaire uniforme donc :

$$\vec{f}_{ie} = m\omega^2\vec{OM} = m\omega^2(x\vec{u}_x + y\vec{u}_y)$$

La force d'inertie de Coriolis est :

$$\vec{f}_{ic} = -2m\vec{\omega} \wedge \vec{v}_{/R'} = -2m\omega\vec{u}_z \wedge (\dot{x}\vec{u}_x + \dot{y}\vec{u}_y) = 2m\omega(\dot{y}\vec{u}_x - \dot{x}\vec{u}_y)$$

On obtient donc en projection sur Ox et Oy :

$$m\ddot{x} = 2m\omega\dot{y} + m\omega^2x \text{ et } m\ddot{y} = -2m\omega\dot{x} + m\omega^2y$$

6. Pour résoudre ces équations il est nécessaire d'utiliser une variable complexe $\underline{C} = x + iy$, avec $\underline{\dot{C}} = \dot{x} + i\dot{y}$ et $\underline{\ddot{C}} = \ddot{x} + i\ddot{y}$. On combine alors les deux équations en multipliant la seconde par i :

$$m(\ddot{x} + i\ddot{y}) = 2m\omega(\dot{y} - i\dot{x}) + m\omega^2(x + iy) \Rightarrow \underline{\ddot{C}} + 2i\omega\underline{\dot{C}} - \omega^2\underline{C} = 0$$

Le discriminant est $\Delta = -4\omega^2 + 4\omega^2 = 0$; on a donc une racine double $-i\omega$ et la solution est :

$$\underline{C} = (\underline{A}t + \underline{B})e^{-i\omega t}$$

À $t = 0$, le point est en O donc $x = y = 0$ et $\underline{C} = 0$ d'où $\underline{B} = 0$. Pour déterminer $\underline{A}$ il faut utiliser la vitesse initiale :

$$\dot{\underline{C}} = \underline{A}e^{-i\omega t}(1 - i\omega t) \Rightarrow \dot{\underline{C}}(0) = \underline{A} = v_0 \Rightarrow \underline{C} = v_0 t e^{-i\omega t}$$

Pour obtenir x et y , il suffit de calculer les parties réelle et imaginaire de $\underline{C}$:

$$x = v_0 t \cos \omega t \text{ et } y = -v_0 t \sin \omega t$$

7. Le point d'impact correspond à $\sqrt{x_d^2 + y_d^2} = d$ soit $d = v_0 t$. On trouve ainsi $x_d = d \cos(\omega d/v_0)$ et $y_d = -d \sin(\omega d/v_0)$. La déviation γ est négative donc elle s'effectue bien sur la droite. Si l'on tire plus fort, ω/v_0 est plus petit et $|\gamma|$ est plus petit donc le point d'impact est plus proche de la cible.

A.5 On étudie la masse m dans le référentiel non galiléen lié au plateau. Dans ce référentiel, les forces sont le poids $m\vec{g}$, la réaction du plateau $\vec{R}$ et la force d'inertie d'entraînement $-m\vec{z}\vec{u}_z$ (il n'y a pas de force d'inertie de Coriolis du fait de l'absence de rotation). La projection du principe fondamental de la dynamique sur l'axe vertical donne $m\ddot{x} = -mg + R + m\omega^2 \cos \omega t$. La masse ne quitte pas le plateau si $\ddot{x} = 0$ et si $R \geq 0$ à tout instant. Cela implique $R = m(g - a\omega^2 \cos \omega t) \geq 0$. Comme le cosinus est compris entre -1 et 1 , cette inégalité est vérifiée à tout instant si $g \geq a\omega^2$ soit $\omega \leq \sqrt{\frac{g}{a}} = \omega_c = 14 \text{ rad.s}^{-1}$.

A.6 1. On étudie la portière modélisée par la masse M au bout d'une tige de masse négligeable à une distance l de la charnière dans le référentiel non galiléen lié à la voiture. La masse M est soumise à son poids, à la réaction de la tige $\vec{R}$ et à la force d'inertie d'entraînement (il n'y a pas de force d'inertie de Coriolis du fait de l'absence de rotation de la voiture). Le principe fondamental de la dynamique s'écrit $m\vec{a} = m\vec{g} + \vec{R} - m\vec{a}_e$. On se place dans le plan horizontal et on utilise les coordonnées polaires dans ce plan centrées sur la charnière de la portière. On projette la relation précédente sur la direction radiale soit $ml\ddot{\theta} = -ma \sin \theta$ ou $\ddot{\theta} + \frac{a}{l} \sin \theta = 0$.

2. En multipliant cette relation par $\dot{\theta}$, on obtient $\frac{d}{dt} \left(\frac{1}{2} \dot{\theta}^2 - \frac{a}{l} \cos \theta \right) = 0$. En intégrant avec la condition que $\dot{\theta} = 0$ pour $\theta = 0$, on en déduit $\dot{\theta}^2 = 2\frac{a}{l} (\cos \theta - 1)$ et la vitesse $v = l\dot{\theta} = \sqrt{2al(\cos \theta - 1)}$. La vitesse pour $\theta = 90^\circ$ vaut $v = \sqrt{-2al}$ (on note que $a < 0$ du fait de la décélération, la racine carrée a donc un sens).

3. Par un raisonnement analogue, on a pour la phase d'accélération $\ddot{\theta} + \frac{a}{l} \sin \theta = 0$ soit $\frac{1}{2} (\dot{\theta}_{\min}^2 - 0) = \frac{a}{l} (1 - \cos \theta_0)$ en intégrant entre la situation initiale où la portière est ouverte $\theta = \theta_0$ sans vitesse initiale et la situation finale où la portière est fermée soit $\theta = 0$ et $\dot{\theta} = \dot{\theta}_{\min}$. Par conséquent, on a $\dot{\theta} = \sqrt{\frac{2a}{l} (1 - \cos \theta_0)}$.

A.7 1. a) Dans le référentiel $\mathcal{R}$ non galiléen les forces sont :

- le poids,
- la réaction de la barre,
- la force d'inertie d'entraînement (ou centrifuge ici),
- la force d'inertie de Coriolis.

Le poids est égal à $\vec{P} = -mg\vec{u}_z$. La réaction est perpendiculaire à la barre soit $\vec{R} = R_y\vec{u}_y + R_z\vec{u}_z$. Le mouvement du point coïncidant dans $\mathcal{R}$ est circulaire uniforme donc l'accélération d'entraînement est $\vec{a}_e = -\omega^2 \vec{OM}$ et la force d'inertie d'entraînement est $\vec{f}_{ie} = m\omega^2 \rho \vec{u}_x$. Enfin, la force d'inertie de Coriolis est donnée par :

$$\vec{f}_{ic} = -2m\omega \wedge \vec{v}'_{/\mathcal{R}'} = -2m\omega \vec{u}_z \wedge \dot{\rho} \vec{u}_x = -2m\dot{\rho} \omega \vec{u}_y$$

- b) La relation fondamentale dans $\mathcal{R}'$ s'écrit :

$$m \vec{a}_{/\mathcal{R}'} = \vec{P} + \vec{R} + \vec{f}_{ie} + \vec{f}_{ic}$$

soit en projection :

$$\begin{cases} m\ddot{\rho} = m\omega^2\rho \\ m\ddot{y} = 0 = R_y - 2m\dot{\rho}\omega \\ m\ddot{z} = 0 = R_z - mg \end{cases}$$

2. On intègre l'équation en ρ : $\ddot{\rho} - \omega^2\rho = 0$ dont la solution est :

$$\rho(t) = Ae^{\omega t} + Be^{-\omega t} \text{ et } \dot{\rho} = A\omega e^{\omega t} - B\omega e^{-\omega t}$$

À $t = 0$ on a $\rho = 0$ et $\dot{\rho} = v_0$, d'où $A + B = 0$ et $\omega(A - B) = v_0$. Après calcul la solution est :

$$\rho(t) = \frac{v_0}{2\omega} (e^{\omega t} + e^{-\omega t}) = \frac{v_0}{\omega} \sinh \omega t$$

3. La vitesse relative, c'est-à-dire dans $\mathcal{R}'$ est $\dot{\rho}\vec{u}_x$. La vitesse d'entraînement est celle dans $\mathcal{R}$ du point P coïncidant à M . Ce point a un mouvement circulaire uniforme dans $\mathcal{R}$, donc sa vitesse est $\vec{\omega} \wedge \vec{OM}$ soit $\vec{v}_e = \rho\omega\vec{u}_y$. On en déduit la vitesse de M dans $\mathcal{R}$:

$$\vec{v}_{/\mathcal{R}} = \vec{v}_{/\mathcal{R}'} + \vec{v}_e = v_0 \cosh \omega t + v_0 \sinh \omega t$$

4. Avec l'expression de $\rho(t)$, on peut écrire que $t_B = \text{asinh} \frac{D\omega}{v_0}$ soit $t_b = 1,57$ s

5. La trajectoire décrite dans $\mathcal{R}$ lorsque M est en butée est un cercle de rayon D . Pour déterminer la trajectoire avant t_B , il faut intégrer la vitesse sachant qu'en coordonnées polaires, avec ici $\vec{u}_\rho = \vec{u}_x$ et $\vec{u}_\theta = \vec{u}_y$ on a : $\vec{v} = \dot{\rho}\vec{u}_x + \rho\dot{\theta}\vec{u}_y$. Avec les équations précédentes on a $\dot{\theta} = \omega$ et $\theta = \omega t$ puisqu'à $t = 0$, $\theta = 0$. L'équation polaire de la trajectoire pour $t < t_B$ est :

$$\rho(\theta) = \frac{v_0}{\omega} \sinh \theta$$

C'est une spirale.

6. Lorsque la masselote est en butée, $\dot{\rho} = 0$. On en déduit avec la relation fondamentale projetée $R_y = 0$ et $R_z = mg$.

A.8 On étudie la masse m dans le référentiel lié au cerceau. La masse m est soumise à son poids, à la réaction du cerceau $\vec{R}$, à la force d'inertie d'entraînement $\vec{f}_{ie} = m\omega^2\vec{AM}$ et à la force d'inertie de Coriolis $\vec{f}_{ic} = -2m\vec{\omega} \wedge \vec{v}_r$. On applique le théorème du moment cinétique appliqué en O . Le moment du poids et de la réaction est nul car la projection du principe fondamental de la dynamique sur la verticale peut s'écrire $m\vec{g} + \vec{R} = \vec{0}$. La force de Coriolis est colinéaire à OM donc son moment par rapport à O est nul. Seule la force d'inertie d'entraînement a donc un moment non nul. Finalement on obtient $mR^2\ddot{\theta} = -m\omega^2R(AM) \sin \frac{\theta}{2}$. Or $AM^2 = 4R^2 \cos^2 \frac{\theta}{2}$ et $\ddot{\theta} + \omega^2 \sin \theta = 0$.

B.1 On étudie la masse m qui est soumise à son poids, à la réaction de la tige et à la tension du ressort en dehors des forces d'inertie si on se place dans le référentiel lié à la tige.

1. L'accélération d'entraînement se limite au terme centrifuge car il n'y a pas de translation et la vitesse de rotation est constante soit $\vec{a}_e = -\omega^2 x \vec{u}_x$.

L'accélération de Coriolis est $\vec{a}_c = 2\vec{\omega} \wedge \vec{v}_r = 2\omega \dot{x} \vec{u}_y$.

2. Le principe fondamental de la dynamique s'écrit $m\vec{a} = m\vec{g} + \vec{T} + \vec{R} - m\vec{a}_e - m\vec{a}_c$. On obtient

$$m\ddot{x} = m\omega^2 x - k \left(x - \frac{x l_0}{\sqrt{x^2 + z_0^2}} \right). \text{ On note que sur les autres directions, on peut en déduire les}$$

$$\text{expressions des composantes de la réaction par } R_y = 2m\omega \dot{x} \text{ et } R_z = mg + k \left(z_0 - \frac{z_0 l_0}{\sqrt{x^2 + z_0^2}} \right).$$

La méthode énergétique est applicable car le poids, la tension du ressort et la force d'inertie d'entraînement (ou force centrifuge) sont des forces conservatives et les autres (réaction et force d'inertie de Coriolis) ne travaillent pas car elles sont perpendiculaires au mouvement.

3. Les positions d'équilibre correspondent à $\ddot{x} = 0$ dans l'expression du principe fondamental de la dynamique ou à l'annulation de la dérivée de l'énergie potentielle pour une résolution énergétique.

$$\text{On obtient } x_e = 0 \text{ et } x_e = \pm \sqrt{\left(\frac{kl_0}{k - m\omega^2}\right)^2 - z_0^2} \text{ si } \omega < \sqrt{\frac{k}{m}}.$$

4. Pour étudier leur stabilité, on étudie le signe de la dérivée seconde de l'énergie potentielle dans le cas d'une résolution énergétique. Sinon on pose $x = x_e + \epsilon$ et on fait un développement limité en ϵ de l'équation du mouvement. On obtient $m\ddot{\epsilon} = \left(\frac{kl_0}{z_0} + m\omega^2 - k\right)\epsilon$. La position d'équilibre est stable si $\frac{kl_0}{z_0} + m\omega^2 - k < 0$. Par conséquent, $x_e = 0$ est une position d'équilibre instable quelle que soit ω si $l_0 < z_0$ et une position d'équilibre stable si $\omega < \sqrt{\frac{k(l_0 - z_0)}{mz_0}}$ sinon.

5. De même, pour les autres positions d'équilibre, on obtient $m\ddot{\epsilon} = -\frac{kl_0x_e^2}{(x_e^2 + z_0^2)^{\frac{3}{2}}}\epsilon$. Elles sont donc stables.

B.2 1. Dans le référentiel galiléen $\mathcal{R}$, les deux forces s'appliquant sur la point P sont le poids et la tension $\vec{T}$ du fil. On applique par exemple le théorème du moment cinétique en O :

$$\frac{d\vec{\sigma}_{/O}}{dt} = \vec{OP} \wedge m\vec{g} + \vec{OP} \wedge \vec{T} \Rightarrow m\ell^2\ddot{\theta} = -mg\ell \sin \theta$$

Le moment de $\vec{T}$ est nul car $\vec{T}$ est colinéaire à $\vec{OM}$. L'équation pour les petites oscillations devient : $\ddot{\theta} + \omega_0^2\theta = 0$ avec $\omega_0 = \sqrt{g/\ell}$. La période est donc $T = 2\pi/\omega_0$. On en déduit :

$$\ell = \frac{T^2g}{4\pi^2} = 0,248 \text{ m}$$

2. La force d'inertie d'entraînement est $\vec{F}_{ie} = -m\vec{a}$ donc son moment par rapport à O' est :

$$\vec{M}'_{O'} = \vec{O'P} \wedge (-m\vec{a}) = -m\ell a \cos \theta \vec{u}_z$$

3. Le référentiel $\mathcal{R}'$ est en translation par rapport à $\mathcal{R}$ donc la force de coriolis est nulle. Son moment par rapport à O' est donc nul aussi.

4. Par rapport au cas où O' est immobile et confondu avec O , il faut ajouter le moment de la force $\vec{F}_{ie}$. Le théorème du moment cinétique dans $\mathcal{R}'$ par rapport à O' (fixe dans $\mathcal{R}'$) s'écrit :

$$\frac{d\vec{\sigma}_{/O'}}{dt} = \vec{O'P} \wedge m\vec{g} + \vec{O'P} \wedge \vec{T} + \vec{O'P} \wedge \vec{F}_{ie}$$

On a alors l'équation suivante projetée sur Oz :

$$m\ell^2\ddot{\theta} = -m\ell(g \sin \theta + a \cos \theta) \Rightarrow \ddot{\theta} = -\frac{g}{\ell} \sin \theta - \frac{a}{\ell} \cos \theta$$

5. La position d'équilibre correspond à $\ddot{\theta} = 0$ donc $\tan \theta_0 = -a/g$.
6. On appelle ϵ l'écart par rapport à la position d'équilibre θ_0 , soit $\theta = \theta_0 + \epsilon$. Il faut effectuer un développement limité au premier ordre de $\cos \theta$ et $\sin \theta$ autour de θ_0 ce qui donne :

$$\cos \theta = \cos(\theta_0 + \epsilon) = \cos \theta_0 \cos \epsilon - \sin \theta_0 \sin \epsilon = \cos \theta_0 - \epsilon \sin \theta_0$$

et

$$\sin \theta = \sin(\theta_0 + \epsilon) = \cos \theta_0 \sin \epsilon + \sin \theta_0 \cos \epsilon = \epsilon \cos \theta_0 + \sin \theta_0$$

Avec $\ddot{\theta} = \ddot{\epsilon}$, l'équation du mouvement devient donc :

$$\ddot{\epsilon} = \frac{-1}{\ell}(g \sin \theta_0 + a \cos \theta_0) - \frac{\epsilon}{\ell}(g \cos \theta_0 - a \sin \theta_0)$$

avec la condition d'équilibre, le premier terme du second membre est nul et le second s'écrit :

$$g \cos \theta_0 - a \sin \theta_0 = g \left(\cos \theta_0 - \frac{a}{g} \sin \theta_0 \right) = g(\cos \theta_0 + \tan \theta_0 \sin \theta_0)$$

ou encore :

$$g \frac{\cos^2 \theta_0 + \sin^2 \theta_0}{\cos \theta_0} = g \sqrt{1 + \tan^2 \theta_0} = \sqrt{g^2 + a^2}$$

L'équation du mouvement devient :

$$\ddot{\epsilon} + \frac{\sqrt{g^2 + a^2}}{\ell} \epsilon = 0 \Rightarrow T = 2\pi \sqrt{\frac{\ell}{\sqrt{g^2 + a^2}}}$$

B.3 1. Puisque l'angle θ sur la figure est égal à ωt , on peut en conclure que $\omega = \dot{\theta}$. En effet, on a $\omega = \dot{\theta}$, donc pour pouvoir intégrer en ωt , il faut ω constant.

Le mouvement du point coïncidant avec M est circulaire uniforme dans $\mathcal{R}$ donc l'accélération d'entraînement (centripète) est $\vec{a}_e = -\omega^2 r \vec{u}_r$. La force d'inertie d'entraînement $\vec{f}_{ie} = -m \vec{a}_e$ est donc $\vec{f}_{ie} = m\omega^2 r \vec{u}_r$.

La force d'inertie de Coriolis est donnée par $\vec{f}_{ic} = -2m\vec{\omega} \wedge \vec{v}_{/\mathcal{R}}$ soit :

$$\vec{f}_{ic} = -2m\omega \vec{r} \vec{u}_z \wedge \vec{u}_r = -2m\omega r \vec{u}_\theta$$

2. Pour calculer l'énergie potentielle dont dérive $\vec{f}_{ie}$, on calcule le travail :

$$\delta W = \vec{f}_{ie} \cdot d\vec{r} = m\omega^2 r dr = d\left(\frac{1}{2}m\omega^2 r^2\right)$$

or $\delta W = -dE_{p,ie}$ d'où $E_{p,ie} = -\frac{1}{2}m\omega^2 r^2$ en choisissant une constante d'intégration nulle avec $E_{p,ie} = 0$ pour $r = 0$.

3. Dans le référentiel $\mathcal{R}'$, la force d'inertie de Coriolis est perpendiculaire à la vitesse de par son expression, sa puissance est nulle.

4. Le mouvement est horizontal, donc le poids n'intervient pas dans le bilan énergétique, pas plus que la réaction normale au support et donc au mouvement. Pour obtenir l'énergie potentielle totale, il faut ajouter à $E_{p,ie}$ l'énergie potentielle dont dérive la force de tension du ressort, soit :

$$E_p = \frac{1}{2}k(r - \ell_0)^2 - \frac{1}{2}m\omega^2 r^2$$

Pour tracer la courbe, on calcule la dérivée :

$$\frac{dE_p}{dr} = k(r - \ell_0) - m\omega^2 r$$

La dérivée peut s'annuler en $r_1 = k\ell_0/(k - m\omega^2)$ si $\omega < \sqrt{k/m}$. On obtient les courbes de la figure (S27.1).

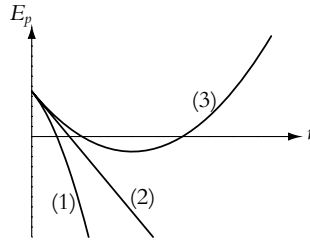


Figure S27.1

Les trois cas sont

- Cas(1) pour $\omega > \sqrt{k/m}$, parabole dans la partie décroissante.
 - Cas(2) pour $\omega = \sqrt{k/m}$, il s'agit d'une droite.
 - Cas(3) pour $\omega < \sqrt{k/m}$, parabole avec minimum en $r = r_1$.
5. Pour avoir une position d'équilibre, il faut un extremum. La seule possibilité est pour $\omega < \sqrt{k/m}$, car pour $\omega > \sqrt{k/m}$ la position $r = 0$ apparaît comme un maximum parce que la zone $r < 0$ n'est pas physique. En fait, la dérivée de E_p est négative pour toute valeur de r .
La longueur ℓ_2 correspondant à l'équilibre est donc $r_1 = \frac{k\ell_0}{k - m\omega^2}$.
6. Comme on l'a vu précédemment, la position d'équilibre n'existe que si $\omega < \sqrt{\frac{k}{m}}$. Elle est stable car elle correspond à un minimum d'énergie potentielle.
7. Si le point M occupe cette position d'équilibre dans $\mathcal{R}'$, il a dans $\mathcal{R}$ un mouvement circulaire uniforme. Mais pour cela, il faut avoir lancé le mobile avec les bonnes conditions initiales, c'est-à-dire avec $\ell_1 = \ell_2$ et avec une vitesse perpendiculaire à $\overrightarrow{OM}$. Dans le cas contraire, il oscille dans $\mathcal{R}'$ autour de la position d'équilibre et donc décrit une sinusoïde variant autour d'un cercle dans $\mathcal{R}$.
8. Avec les conditions initiales choisies, on a bien la vitesse initiale perpendiculaire à $\overrightarrow{OM}$. Il suffit donc de choisir $\ell_1 = \ell_2$ pour que M soit à la position d'équilibre.

B.4 1. Le référentiel $\mathcal{R}_S$ est en rotation par rapport au référentiel galiléen du laboratoire. Les forces qui s'exercent sur un point M dans $\mathcal{R}_S$ sont le poids (conservative), la réaction du bol (ne travaille pas), la force d'inertie d'entraînement (conservative) et la force de coriolis (ne travaille pas). Le mouvement par rapport à $\mathcal{R}_L$ du point P coïncidant à M est un mouvement circulaire uniforme d'axe Oz et de vitesse angulaire ω , donc l'accélération d'entraînement est $\vec{a}_e = -\omega^2 \overrightarrow{HM}$ où H est la projection orthogonale de M sur l'axe Oz . En coordonnées cylindriques, la force d'inertie d'entraînement $\vec{f}_{ie} = -m\vec{a}_e$ est donc égale à $m\omega^2 \rho \vec{u}_\rho$. Pour calculer l'énergie potentielle dont elle dérive on calcule le travail :

$$\delta W = \vec{f}_{ie} \cdot d\vec{\rho} = m\omega^2 \rho d\rho = -dE_{p,ie} \Rightarrow E_{p,ie} = -\frac{1}{2}m\omega^2 \rho^2 + cte$$

L'énergie potentielle pour le poids est $mgz + cte$ donc finalement :

$$E_p = mgz - \frac{1}{2}m\omega^2 \rho^2 + cte$$

Pour un point qui est sur la sphère $\rho^2 + z^2 = r_0^2$, on peut donc écrire :

$$E_p = mgz - \frac{1}{2}m\omega^2 (r_0^2 - z^2) + cte = \frac{1}{2}m\omega^2 z \left(z + \frac{2g}{\omega^2} \right) - \frac{1}{2}m\omega^2 r_0^2 + cte$$

Ce qui donne finalement l'expression de l'énoncé si on choisit la constante telle que $E_p(z = 0) = 0$.

2. On calcule la dérivée de l'énergie potentielle :

$$\frac{dE_p}{dz} = \frac{1}{2}m\omega^2 \left(\frac{\omega_0^2}{\omega^2} r_0 + 2z \right)$$

On a donc un extremum pour $z_e = -\frac{r_0}{2} \frac{\omega_0^2}{\omega^2}$. Le signe négatif correspond bien à la zone du bol (en-dessous de O). Pour que la position existe, il faut $|z_e| \leq r_0$ donc $\omega \geq \omega_0/\sqrt{2}$. Pour la stabilité, on calcule la dérivée seconde :

$$\frac{d^2E_p}{dz^2} = m\omega^2 > 0$$

donc la position est stable si elle existe.

La courbe z_e fonction de ω/ω_0 est représentée sur la figure (S27.2). On remarque que lorsque ω croît, z_e se rapproche de 0 et la position d'équilibre se rapproche du bord du bol en s'éloignant de l'axe : c'est l'effet de la force d'inertie qui est une force centrifuge.

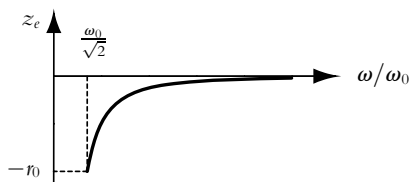


Figure S27.2

3. a) Le courbe $E_p(z)$ est représentée sur la figure (S27.3).

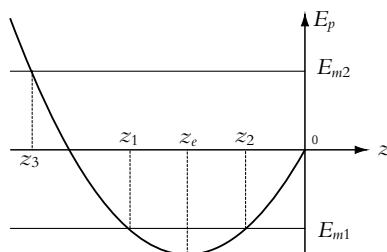


Figure S27.3

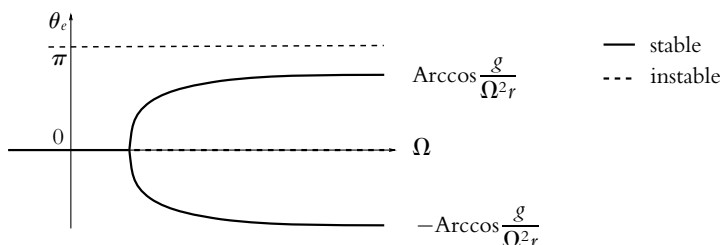
- b) La courbe d'énergie potentielle montre que l'on a un puits de potentiel. Le mouvement peut se faire dans les zones où $E_m \geq E_p$ puisque l'énergie cinétique est une grandeur positive ou nulle. Si on lâche la bille à $z_0 = z_e$, alors la seule possibilité est qu'elle reste à cette position. Elle va donc tourner avec le bol. Pratiquement, ce cas est impossible à réaliser car il faut avoir exactement $z_0 = z_e$. Si on lâche la bille à une autre cote, pour laquelle l'énergie mécanique E_{m1} est négative, elle a une énergie mécanique supérieure à $E_{p,min}$ elle va alors osciller entre deux cotes z_1 et z_2 (figure S27.3). Si on lâche la bille à une cote, pour laquelle l'énergie mécanique E_{m2} est positive elle va osciller d'un côté à l'autre en passant par le fond, et en remontant à la cote de départ z_3 de chaque côté.

B.5 1. On adopte comme d'habitude la méthode de résolution en quatre points :

- Système : point matériel A .
- Référentiel $\mathcal{R}$ non galiléen lié à la piste et en rotation autour de son diamètre vertical.
- Bilan des forces :
 - poids $m\vec{g}$,
 - réaction du support $\vec{R}$ perpendiculaire à la piste du fait de l'absence de frottement,

- force d'inertie d'entraînement $\vec{f}_{ie} = m\Omega^2 \vec{HM}$ (en notant H la projection orthogonale de M sur l'axe de rotation),
 - force d'inertie de Coriolis : $\vec{f}_{ic} = -2m\vec{\Omega} \wedge \vec{v}/_{\mathcal{R}}(M)$.
2. La réaction et la force d'inertie de Coriolis sont perpendiculaires à la trajectoire : ces deux forces ne travaillent pas. Le poids et la force d'inertie d'entraînement (qui se limite ici à la force centrifuge) dérivent d'un potentiel. Enfin la position du point A est déterminée par un seul paramètre, l'angle θ . Une méthode énergétique est donc adaptée puisque les forces sont conservatives ou ne travaillent pas et que le problème n'a qu'un paramètre de position.
3. On utilise les coordonnées polaires dans le plan de la piste. La vitesse s'exprime alors par : $\vec{v} = r\dot{\theta}\vec{u}_\theta$. L'énergie cinétique est donc : $E_c = \frac{1}{2}mv^2 = \frac{1}{2}mr^2\dot{\theta}^2$.
4. La différentielle de l'énergie potentielle de pesanteur s'écrit :
 $dE_{\text{pesanteur}} = -m\vec{g} \cdot d\vec{OM} = -mgr \sin \theta d\theta$ donc $E_{\text{pesanteur}} = -mgr \cos \theta + C$ où C est une constante. De même, la différentielle de l'énergie potentielle de la force centrifuge s'écrit :
 $dE_{\text{centrifuge}} = -m\Omega^2 \vec{HM} \cdot d\vec{OM} = -m\Omega^2 r \sin \theta r \cos \theta d\theta$ donc $E_{\text{centrifuge}} = -\frac{1}{2}m\Omega^2 r^2 \sin^2 \theta + C'$ où C' est une constante.
 Soit finalement : $E_p = E_{\text{pesanteur}} + E_{\text{centrifuge}} = -mgr \cos \theta - \frac{1}{2}m\Omega^2 r^2 \sin^2 \theta + K$ où K est une constante qu'on détermine en choisissant l'origine des potentiels en $\theta = \frac{\pi}{2}$ c'est-à-dire :
 $-\frac{1}{2}m\Omega^2 r^2 + K = 0$ et finalement : $E_p = -mr \cos \theta \left(g - \frac{\Omega^2 r \cos \theta}{2} \right)$.
5. Comme toutes les forces sont conservatives ou ne travaillent pas, l'énergie mécanique $E_m = E_c + E_p$ se conserve. Par dérivation par rapport au temps, on obtient : $mr^2\dot{\theta}\ddot{\theta} + mgr \sin \theta \dot{\theta} - m\Omega^2 r^2 \cos \theta \sin \theta \dot{\theta} = 0$ soit $\dot{\theta} = 0$ ce qui correspond à l'absence de mouvement ou $\ddot{\theta} + \sin \theta \left(\frac{g}{r} - \Omega^2 \cos \theta \right) = 0$ qui est l'équation du mouvement.
6. Les positions d'équilibre correspondent aux extrema de l'énergie potentielle ou aux positions correspondant à $\ddot{\theta} = \dot{\theta} = 0$. Dans les deux cas, on obtient : $\sin \theta \left(\frac{g}{r} - \Omega^2 \cos \theta \right) = 0$ c'est-à-dire $\sin \theta = 0$ ou $\cos \theta = \frac{g}{\Omega^2 r}$, cette position n'existant que si $g < \Omega^2 r$. Les positions d'équilibre sont donc : $\theta = 0$, $\theta = \pi$ et, si $\frac{g}{\Omega^2 r} \leq 1$, $\theta = \pm \text{Arccos} \frac{g}{\Omega^2 r}$.
7. Les positions d'équilibre sont stables si la dérivée seconde de l'énergie potentielle est positive. Or $\frac{d^2 E_p}{d\theta^2} = mgr \cos \theta + m\Omega^2 r^2 (1 - 2 \cos^2 \theta)$.
- En $\theta = 0$, $\frac{d^2 E_p}{d\theta^2} = mgr - m\Omega^2 r^2 > 0$ si $g \leq \Omega^2 r$ (on note que cette condition est également celle de l'inexistence des positions d'équilibre $\pm \text{Arccos} \frac{g}{\Omega^2 r}$).
 - En $\theta = \pi$, $\frac{d^2 E_p}{d\theta^2} = -mgr - m\Omega^2 r^2 < 0$.
 - En $\theta = \pm \text{Arccos} \frac{g}{\Omega^2 r}$, $\frac{d^2 E_p}{d\theta^2} = m \frac{g^2}{\Omega^2} + m\Omega^2 r^2 \left(1 - 2 \frac{g^2}{\Omega^4 r^2} \right) = -m \frac{g^2}{\Omega^2} + m\Omega^2 r^2 > 0$ si $g \geq \Omega^2 r$ (il s'agit également de la condition d'existence de ces positions d'équilibre).
- Cette étude permet d'obtenir les conclusions suivantes :
- Si $\Omega^2 r \leq g$ alors on a deux positions d'équilibre : l'une stable $\theta = 0$ et l'autre instable $\theta = \pi$.
 - Si $\Omega^2 r \geq g$ alors on a quatre positions d'équilibre : deux stables $\theta = \pm \text{Arccos} \frac{g}{\Omega^2 r}$ et deux instables $\theta = 0$ et $\theta = \pi$.

8.



B.6 1. À l'équilibre dans $\mathcal{R}$, les forces qui s'exercent sur M sont :

- Le poids : $\vec{P} = -mg(\cos \alpha \vec{u}_X + \sin \alpha \vec{u}_Y)$,
- la réaction normale de l'axe : $\vec{R} = R\vec{u}_Y$,
- la tension du ressort : $\vec{T} = -k(X - \ell_0)\vec{u}_X$.

La relation fondamentale à l'équilibre en projection sur PX entraîne :

$$-mg \cos \alpha - k(X_e - \ell_0) = 0 \Rightarrow X_e = \ell_0 - \frac{mg \cos \alpha}{k}$$

2. Le référentiel $\mathcal{R}'$ est en translation par rapport à $\mathcal{R}$ donc la force d'inertie de coriolis est nulle. L'accélération d'entraînement est donc l'accélération de P dans $\mathcal{R}$ soit :

$$\vec{a}_e = \frac{d\vec{OP}}{dt} = -a\omega^2 \cos \omega t \vec{u}_z \Rightarrow \vec{f}_{ie} = ma\omega^2 \cos \omega t \vec{u}_z$$

soit en projection sur les vecteurs liés à la barre :

$$\vec{f}_{ie} = ma\omega^2 \cos \omega t (\cos \alpha \vec{u}_X + \sin \alpha \vec{u}_Y)$$

3. La relation fondamentale de la dynamique dans $\mathcal{R}'$ est :

$$m \vec{a}_{/\mathcal{R}'} = m \vec{g} + \vec{R} + \vec{T} + \vec{F}_f + \vec{f}_{ie}$$

soit en projection sur les axes :

$$\begin{cases} m\ddot{X} = -mg \cos \alpha - k(X - \ell_0) - \lambda \dot{X} + ma\omega^2 \cos \alpha \cos \omega t \\ m\ddot{Y} = 0 = -mg \sin \alpha + R + ma\omega^2 \sin \alpha \sin \omega t \end{cases}$$

4. L'équation en X peut s'écrire :

$$\ddot{X} + \frac{\lambda}{m} \dot{X} + \frac{k}{m} (X - X_e) = a\omega^2 \cos \alpha \cos \omega t$$

Comme X_e est une constante, alors $\ddot{E} = \ddot{X}$ et $\dot{E} = \dot{X}$. On trouve l'équation demandée avec $\omega_0^2 = k/m$, $Q = m\omega_0/\lambda$ et $\alpha = a \cos \alpha$.

5. En complexe l'équation s'écrit :

$$\underline{E} \left(-\omega^2 + i \frac{\omega \omega_0}{Q} + \omega_0^2 \right) = \alpha \omega^2 e^{i\omega t}$$

on en tire le module de $\underline{E}$:

$$E = \frac{\alpha \omega^2}{\sqrt{(\omega^2 - \omega_0^2)^2 + \left(\frac{\omega \omega_0}{Q} \right)^2}}$$

6. Le changement de variable conduit à :

$$E = \frac{\alpha}{\sqrt{(1 - \Sigma^2)^2 + \left(\frac{\Sigma}{Q} \right)^2}}$$

On peut écrire E sous la forme $E = \alpha / \sqrt{f(\Sigma)}$ soit :

$$\frac{dE}{d\Sigma} = -\frac{1}{2} \frac{\alpha}{f^{3/2}} \frac{df}{d\Sigma}$$

alors :

$$\frac{df}{d\Sigma} = 4\Sigma(\Sigma^2 - 1) + 2\frac{\Sigma}{Q^2} = 4\Sigma \left(\Sigma^2 - 1 + \frac{1}{2Q^2} \right)$$

La dérivée est nulle pour $\Sigma = 0$ et pour $\Sigma_r = \sqrt{1 - \frac{1}{2Q^2}}$ si $Q > 1/\sqrt{2}$.

Il se présente donc deux cas :

- Si $Q < 1/\sqrt{2}$, alors la dérivée de f est positive pour tout Σ , donc celle de E est négative : E est une fonction décroissante de Σ donc croissante de ω .
- Si $Q > 1/\sqrt{2}$, alors la dérivée de f est négative pour $\Sigma < \Sigma_r$ et positive pour $\Sigma > \Sigma_r$. Ainsi la dérivée de E est positive pour $\Sigma < \Sigma_r$ et négative pour $\Sigma > \Sigma_r$; E admet donc un maximum (résonance) pour Σ_r ou pour $\omega_r = \omega_0/\sqrt{1 - \frac{1}{2Q^2}}$.

Les courbes représentatives sont tracées sur la figure (S27.4). La courbe (2) pour $Q > 1/\sqrt{2}$ et la courbe (1) pour $Q < 1/\sqrt{2}$.

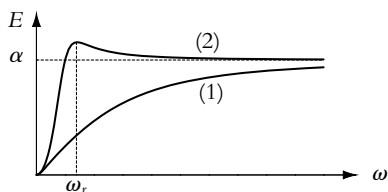


Figure S27.4

Ce filtre est passe-haut car $E(\omega = 0) = 0$ et $E(\omega \rightarrow \infty) = \alpha$ donc il coupe les basses fréquences et laisse passer les hautes fréquences.

7. Pour que le système donne une amplitude proportionnelle à a , il faut procéder aux hautes fréquences $\omega > \omega_r$ et l'amplitude vaut alors α .
Sans le système de freinage, le régime transitoire est infiniment long.

Chapitre 29

A.1 1. Faux. Le théorème de l'énergie cinétique dit que la variation d'énergie cinétique est égale à la somme des travaux des forces intérieures et extérieures. Si le système est isolé, il n'y a pas de forces extérieures. L'énergie mécanique est la somme de l'énergie cinétique et des énergies potentielles. Dans le cas d'un système isolé, il faut donc pouvoir écrire tous les travaux des forces intérieures sous la forme d'une énergie potentielle, ce qui n'est pas toujours le cas. Le fait que le système soit isolé n'est donc pas suffisant pour assurer la conservation de l'énergie mécanique.

2. Vrai. Le fait que le travail des forces ne dépende pas du chemin suivi signifie qu'on peut l'écrire comme la différentielle d'une fonction qui sera l'opposé de l'énergie potentielle. Il s'agit d'une condition assurant la conservation de l'énergie mécanique.

3. Vrai. Il s'agit d'une autre formulation de la condition de la question précédente.

4. Faux. L'énergie cinétique peut être constante sans que l'énergie mécanique soit conservée : il suffit que le travail d'une force conservative compense celui d'une force non conservative par exemple.

A.2 1. On appelle $\mathcal{R}$ le référentiel galiléen lié à l'axe Ox .

La vitesse du centre de gravité dans $\mathcal{R}$ juste après le choc est :

$$\vec{V}_{G,0} = \frac{m_1 \vec{v}_{1,0} + m_2 \vec{v}_{2,0}}{m_1 + m_2} = \frac{m_1}{m_1 + m_2} \vec{v}_0$$

Les forces extérieures s'exerçant sur le système global (M_1 , M_2 , ressort), sont les poids des deux points et les réactions normales du support (perpendiculaire à l'axe Ox). La projection du théorème du centre d'inertie sur Ox donne :

$$\frac{d\vec{v}_G}{dt} = m_1 \vec{g} + m_2 \vec{g} + \vec{R}_{N1} + \vec{R}_{N2} \Rightarrow \ddot{x}_G = 0 \Rightarrow \dot{x}_G = ct = \frac{m_1}{m_1 + m_2} v_0$$

2. a) Dans le référentiel barycentrique

Puisque G a un mouvement rectiligne uniforme, le référentiel barycentrique $\mathcal{R}_B$ est galiléen. On définit la particule fictive M par $\vec{GM} = \vec{M_1M_2}$, et l'on suppose que $\vec{M_1M_2} = (M_1M_2)\vec{u}_x$. On pose aussi $\vec{GM} = x\vec{u}_x$.

On a vu dans le cours que les forces s'exerçant sur M sont alors celles s'exerçant sur M_2 , c'est-à-dire :

- le poids $-m_2g\vec{u}_z$,
- la réaction normale $\vec{R}_{N2}$,
- la tension du ressort $\vec{T}_2 = -k(\ell - \ell_0)\vec{u}_x$,
où la longueur ℓ du ressort est égale à M_1M_2 donc à x . La relation fondamentale de la dynamique appliquée à M dans $\mathcal{R}_B$ s'écrit :

$$\mu \frac{d^2x}{dt^2} = -m_2g\vec{u}_z + \vec{R}_{N2} + \vec{T}_2 \Rightarrow \mu \ddot{x} = -k(x - \ell_0) \Rightarrow \ddot{x} + \omega_0^2 x = \omega_0^2 \ell_0$$

avec $\omega_0 = \sqrt{k/\mu}$ où μ est la masse réduite.

La solution est donc $x(t) = A \cos \omega_0 t + B \sin \omega_0 t + \ell_0$. À $t = 0$ la longueur est égale à ℓ_0 soit $\ell_0 = A + \ell_0$, donc $A = 0$. On calcule la vitesse : $\dot{x} = B\omega_0 \cos \omega_0 t$. À $t = 0$, la vitesse de M est égale à $v_{2,0} - v_{1,0} = -v_0$. On en déduit $B = -v_0/\omega_0$. Finalement :

$$\ell(t) = x(t) = -v_0 \sqrt{\frac{\mu}{k}} \sin \sqrt{\frac{k}{\mu}} t + \ell_0$$

b) Dans le référentiel absolu

La longueur du ressort est $\ell = x_2 - x_1$. On projette la relation fondamentale de la dynamique appliquée à chaque particule sur Ox en faisant attention aux sens opposés des deux tensions :

$$\begin{cases} m_1 \ddot{x}_1 = k(\ell - \ell_0) \\ m_2 \ddot{x}_2 = -k(\ell - \ell_0) \end{cases}$$

On divise la deuxième équation par m_2 et on soustrait la première divisée par m_1 ce qui donne :

$$\ddot{x}_2 - \ddot{x}_1 = -\left(\frac{1}{m_2} + \frac{1}{m_1}\right) k(\ell - \ell_0) \Rightarrow \ddot{\ell} = -\frac{k}{\mu}(\ell - \ell_0)$$

On retrouve alors l'équation résolue dans la question précédente avec les mêmes conditions initiales puisque $\dot{\ell}(t=0) = \dot{x}_{2,0} - \dot{x}_{1,0}$.

B.1 On étudie dans un premier temps le système formé des deux masses m_1 et m_2 et du fil inextensible et on se place dans le référentiel du laboratoire considéré comme galiléen. Les forces extérieures sont le poids de chacune des masses et la réaction du support sur la masse m_1 qu'on décompose en une composante normale R_N et une composante tangentielle R_T . Les forces intérieures sont les tensions du fil sur les masses et celles des masses sur le fil.

Si on applique le théorème de l'énergie cinétique, il faut calculer le travail de toutes les forces. Cependant on peut remarquer que le travail de la tension du fil sur une masse et celui de la tension de la masse sur le fil sont opposés et par conséquent s'annulent. Le poids de la masse m_1 et la composante normale R_N de la réaction du support sur m_1 ne travaillent pas puisque ces forces sont perpendiculaires au déplacement. Le travail du poids de la masse m_2 vaut : $W_2 = m_2gh$.

La projection du principe fondamental de la dynamique appliqué à la masse m_1 sur la verticale où il n'y a pas de mouvement donne : $R_N = m_1g$. On en déduit par la loi de Coulomb que $R_T = f m_1g$. Puisque R_T

est constante, on en déduit que le travail de la force de frottement solide vaut : $W = -hfm_1g$. Finalement le théorème de l'énergie cinétique appliqué à l'ensemble donne : $\frac{1}{2}(m_1 + m_2)v^2 - 0 = m_2gh - fm_1gh$.

Quand la masse m_2 atteint le sol, le fil se détend et il n'y a plus de tension du fil. L'énergie cinétique de la masse m_1 va alors diminuer à cause de la force de frottement. On étudie maintenant le système constitué de la masse m_1 seule. Elle est soumise à son poids et à la réaction du support qu'on décompose toujours en un composante normale R_N et une composante tangentielle R_T . Comme pour la première phase du mouvement, la projection du principe fondamental de la dynamique appliqué à la masse m_1 sur la verticale donne : $R_N = m_1g$. D'autre part, R_N et le poids de la masse m_1 ne travaillent pas puisqu'elles sont perpendiculaires au déplacement. Il ne reste donc que le travail de la force de frottement qui vaut $W = -fm_1gd$. Alors $0 - \frac{1}{2}m_1v^2 = -fm_1gd$ par application du théorème de l'énergie cinétique à la seule masse m_1 pendant le trajet de longueur d .

Les deux expressions de v^2 donnent : $v^2 = 2\frac{(m_2 - fm_1)gh}{m_1 + m_2} = 2fgd$, on en déduit la valeur du coefficient de frottement : $f = \frac{m_2h}{(m_1 + m_2)d + m_1h}$.

B.2 1. a) On considère le système $(B + C)$. Dans le référentiel galiléen $\mathcal{R}$ lié au sol, les forces extérieures s'exerçant sur le système sont le poids de C , le poids de B et la réaction normale du support en C . On note G le centre de gravité du système. Comme ces trois forces sont verticales, le théorème du centre d'inertie donne :

$$\frac{d\vec{v}_G}{dt} = m\vec{g} + M\vec{g} + \vec{R}_n \Rightarrow \ddot{x}_G = 0 \Rightarrow \dot{x}_G = ct$$

or à $t = 0$ le système est immobile donc $\dot{x}_G = 0$. D'autre part, on peut écrire :

$$\dot{x}_G = \frac{M\dot{x}_C + m\dot{x}_B}{M + m}$$

Pour calculer la vitesse de B dans $\mathcal{R}$, il faut utiliser la composition des vitesses. Le mouvement de B par rapport à A se fait sur un cercle de centre A à vitesse angulaire $\dot{\theta}$, donc en projection sur Ox :

$$\dot{x}_B = \dot{x}_C + L\dot{\theta} \cos \theta$$

En combinant les différentes relations, on obtient :

$$\dot{x} + \alpha L\dot{\theta} \cos \theta = 0$$

b) Puisque la vitesse horizontale $\dot{x}_G$ de G est nulle, l'abscisse x_G est constante et G se déplace sur une verticale.

c) On utilise l'une des relations de définition de G que l'on projette sur Ox :

$$M\vec{GC} + m\vec{GB} = \vec{0} \Rightarrow MX_C + mX_B = 0$$

La relation précédente montre que X_B et X_C sont de signes opposés. Les deux points oscillent de part et d'autre du point G , lorsque B est à droite, C est à gauche et réciproquement.

Pour exprimer la cote de G , on utilise l'autre relation de définition de G que l'on projette sur Oz :

$$(m + M)\vec{OG} = M\vec{OC} + m\vec{OB} \Rightarrow z_G = \frac{m}{M + m}(h - L \cos \theta) = \alpha(h - L \cos \theta)$$

d) La projection de la relation de définition de G sur Ox donne :

$$(m + M)x_G = Mx_C + mx_B = Mx_C + m(x_C + L \sin \theta) \Rightarrow (m + M)(x_G - x_C) = mL \sin \theta$$

or $x_C - x_G = X_C$, d'où :

$$X_B = -\frac{M}{m}X_C = \frac{M}{M + m}L \sin \theta = (1 - \alpha)L \sin \theta$$

Pour déterminer Z_B il suffit d'écrire :

$$Z_B = z_B - z_G = (1 - \alpha)(h - L \cos \theta)$$

On peut alors écrire l'équation de la trajectoire de B :

$$\left(\frac{X_B}{(1-\alpha)L}\right)^2 + \left(\frac{1}{L}\left(\frac{Z_B}{(1-\alpha)} - h\right)\right)^2 = \sin^2 \theta + \cos^2 \theta = 1$$

Le point B se déplace sur une ellipse.

2. Lorsque le chariot est contre la butée, celle-ci exerce une force de réaction horizontale, donc $\ddot{x}_G$ n'est plus nulle. Dès que le chariot ne touche plus la butée, $\ddot{x}_G$ redevient nulle. La différence avec le cas précédent réside dans le fait que la vitesse n'est plus nulle. Si on appelle v_{0x} la vitesse horizontale de G lorsque le chariot se détache de la butée, on a alors :

$$v_{0x} = \dot{x} + \alpha L \dot{\theta} \cos \theta$$

tant que le chariot ne sera pas en butée.

3. a) L'énergie cinétique totale est la somme de l'énergie cinétique de B et de C :

$$E_c = \frac{1}{2} M \dot{x}^2 + \frac{1}{2} m v_B^2$$

Si l'on note $\vec{u}_\theta$, le vecteur orthonormal des coordonnées polaires repérant B par rapport à A , la composition des vitesses entraîne :

$$\vec{v}_B = \vec{v}_A + L \dot{\theta} \vec{u}_\theta = \dot{x} \vec{u}_x + L \dot{\theta} \vec{u}_\theta \Rightarrow v_B^2 = \dot{x}^2 + L^2 \dot{\theta}^2 + 2 \dot{x} L \dot{\theta} \cos \theta$$

En remplaçant $\dot{x}$ par son expression en fonction de θ , on obtient :

$$E_c = \frac{1}{2} [(M+m)\alpha^2 - 2m\alpha] L^2 \dot{\theta}^2 \cos^2 \theta + mL^2 \dot{\theta}^2$$

- b) L'énergie potentielle est l'énergie potentielle de pesanteur. Si l'on prend l'origine en $z = 0$, seule l'énergie de B intervient, soit :

$$E_p = mgz_B = mg(h - L \cos \theta)$$

- c) L'énergie mécanique totale est $E_m = E_c + E_p$. Elle est constante au cours du mouvement puisqu'on a affaire à deux forces conservatives (les poids) et une réaction normale, perpendiculaire au mouvement de C et donc qui ne travaille pas. Pour obtenir l'équation du mouvement, on exprime le fait que la dérivée de E_m par rapport au temps est nulle :

$$\frac{dE_m}{dt} = \dot{\theta} [(M+m)\alpha^2 - 2m\alpha] L^2 (\ddot{\theta} \cos^2 \theta - \dot{\theta}^2 \cos \theta \sin \theta) + mL^2 \ddot{\theta} + mgL \sin \theta$$

On exclut l'immobilité, donc on peut simplifier par le $\dot{\theta}$ en facteur et on obtient l'équation du mouvement :

$$((M+m)\alpha^2 - 2m\alpha) L^2 (\ddot{\theta} \cos^2 \theta - \dot{\theta}^2 \cos \theta \sin \theta) + mL^2 \ddot{\theta} + mgL \sin \theta = 0$$

Pour obtenir l'équation des petites oscillations, on effectue un développement limité au premier ordre et on ne garde que les termes de premier ordre, c'est-à-dire qu'on néglige les termes en θ^2 et $\dot{\theta}^2$. On a alors $\sin \theta \simeq \theta$ et $\cos \theta \simeq 1$. L'équation obtenue est :

$$((M+m)\alpha^2 - 2m\alpha + m) L^2 \ddot{\theta} + mgL \theta = 0$$

La période des petites oscillations est :

$$T_0 = 2\pi \left(\frac{((M+m)\alpha^2 - 2m\alpha + m)L}{mg} \right)^{1/2}$$

- B.3** 1. Le référentiel du laboratoire est galiléen, et le système considéré est $(A, B, \text{ressort})$. Les forces extérieures sont : les poids, les forces de réaction normales du support et la tension du fil $\vec{T}_f$. Les forces intérieures sont les tensions sur A et sur B du ressort.

On applique la relation fondamentale de la dynamique au point A en supposant que les trois points ne sont pas alignés (figure S29.1) :

$$m \vec{a}_A = m \vec{g} + \vec{R}_A + \vec{F}_f + \vec{T}_A$$

où $\vec{T}_A$ est la tension exercée par le ressort sur A . La projection de cette équation sur $\vec{u}_r$ et $\vec{u}_\theta$ donne :

$$\begin{cases} m(\ddot{L} - L\dot{\theta}^2) = -T_f + T_{A,r} \\ m(2\dot{L}\dot{\theta} + L\ddot{\theta}) = T_{A,\theta} \end{cases} \Rightarrow \begin{cases} -mL\dot{\theta}^2 = -T_f + T_{A,r} \\ 0 = T_{A,\theta} \end{cases}$$

puisque $L = cte$ et $\omega = \dot{\theta} = cte$. La composante suivant $\vec{u}_\theta$ de la tension du ressort est nulle. Comme la tension est colinéaire au ressort, les trois points sont alignés.

Le point G est à mi-distance de A et B , il est à distance constante de O , donc son mouvement est circulaire de rayon $L + \ell_1/2$. On applique le théorème du centre d'inertie à (A, B) et on projette sur $\vec{u}_r$:

$$2m \frac{d\vec{v}_G}{dt} = \vec{R}_{ext} \Rightarrow -2m \left(L + \frac{\ell_1}{2} \right) \omega^2 = -T_f$$

d'où T_f .

Pour calculer ℓ_1 on applique la relation fondamentale à A , soit en projection sur $\vec{u}_r$:

$$-mL\omega^2 = -T_f + k(\ell_1 - \ell_0)$$

En combinant les deux équations, on obtient : $\ell_1 = \frac{mL\omega^2 + k\ell_0}{k - m\omega^2}$. Dans le cas $k = 2m\omega^2$, on trouve $\ell_1 = L + 2\ell_0$ et $T_f = m\omega^2(2\ell_0 + 3L)$.

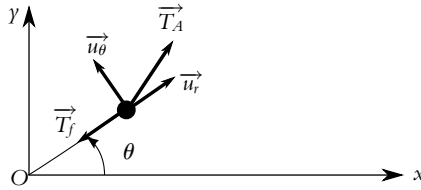


Figure S29.1

2. a) Lorsque la liaison OA n'existe plus, le système est pseudo-isolé (les forces extérieures verticales se compensent et il n'y a pas de forces extérieures horizontales). Le théorème du centre d'inertie donne :

$$2m \frac{d\vec{v}_G}{dt} = \vec{R}_{ext} = \vec{0} \Rightarrow \vec{v}_G = d\vec{e} = \vec{v}_G(0) = \frac{m\vec{v}_A(0) + m\vec{v}_B(0)}{2m}$$

soit finalement

$$\vec{v}_G = \frac{L\omega + (L + \ell_1)\omega}{2} \vec{u}_y$$

Le mouvement de G est donc rectiligne uniforme donc $\vec{OG} = \vec{OG}_0 + \vec{v}_G t$. Dans le cas $k = 2m\omega^2$, on a :

$$\vec{v}_G = \frac{\omega}{2} (3L + 2\ell_0) \vec{u}_y$$

- b) Le référentiel barycentrique $\mathcal{R}_B$ est le référentiel lié à G et en translation (à la vitesse $\vec{v}_G$) par rapport au référentiel absolu. Le mouvement de G est une translation rectiligne uniforme dans $\mathcal{R}$, donc $\mathcal{R}_B$ est aussi un référentiel galiléen.
3. a) On définit $\vec{r} = \vec{r}^* = \vec{AB}$, l'exposant « $*$ » correspondant à l'expression dans $\mathcal{R}_B$. On a alors $\vec{r}_A^* = \vec{GA}$ et $\vec{r}_B^* = \vec{GB}$. La définition de G permet d'écrire $m_A \vec{r}_A^* + m_B \vec{r}_B^* = \vec{0}$. On trouve alors :

$$\vec{r}_B^* = -\vec{r}_A^* = \frac{\vec{r}^*}{2} \text{ et } \vec{p}_B^* = -\vec{p}_A^* = \frac{m}{2} \frac{d\vec{r}^*}{dt}$$

- b) L'énergie cinétique totale est :

$$E_c^* = \frac{1}{2} m_A \vec{v}_A^{*2} + \frac{1}{2} m_B \vec{v}_B^{*2} = \frac{m}{4} \left(\frac{d\vec{r}^*}{dt} \right)^2$$

L'énergie potentielle d'interaction se calcule à partir du travail de la tension aux points A et B :

$$\delta W = \vec{T}_A \cdot d\vec{r}_A^* + \vec{T}_B \cdot d\vec{r}_B^* = k(r^* - \ell_0) d\left(-\frac{r^*}{2}\right) - k(r^* - \ell_0) d\left(\frac{r^*}{2}\right)$$

On peut alors écrire :

$$\delta W = -dE_p = -d\left(\frac{1}{2}k(r^* - \ell_0)^2 + cte\right) \Rightarrow E_p = \frac{1}{2}k(r^* - \ell_0)^2 + cte$$

Le moment cinétique total est :

$$\vec{\sigma}^* = \vec{GA} \wedge m\vec{v}_A^* + \vec{GB} \wedge m\vec{v}_B^* = \vec{r}^* \wedge \frac{m}{2} \frac{d\vec{r}^*}{dt}$$

Les précédentes relations correspondent à celles d'un point M tel que $\vec{GM} = \vec{r}^*$, de vitesse $\frac{d\vec{r}^*}{dt}$ et de masse $\mu = m_A m_B / (m_A + m_B)$ soit ici $\mu = m/2$.

- c) La vitesse initiale de M est égale à $\vec{v}_B(0) - \vec{v}_1(0)$ soit $\vec{V}_1 = \ell_1 \omega \vec{u}_y$. Dans le cas $k = 2m\omega^2$, la vitesse initiale est $\vec{V}_1 = (L + 2\ell_0) \vec{u}_y$.
- d) Se référer au cours pour montrer que la force totale subie par M est celle subie par B soit le poids $m\vec{g}$, la réaction du plan sur B $\vec{R}_B$ et la tension du ressort $\vec{T}_B$ (notée après $\vec{T}$).
- e) Les forces verticales $m\vec{g}$ et $\vec{R}_B$ se compensent donc il ne reste que la tension du ressort qui est colinéaire à $\vec{GM}$. Son moment $\vec{GM} \wedge \vec{T}$ par rapport à G est nul. Le théorème du moment cinétique dans $\mathcal{R}_B$ permet de conclure que le moment cinétique $\vec{\sigma}^*$ est constant :

$$\frac{d\vec{\sigma}^*}{dt} = \vec{GM} \wedge \vec{T} = \vec{0}$$

D'autre part, la réaction et le poids perpendiculaires au mouvement ne travaillent pas et la tension du ressort est une force conservative donc l'énergie mécanique E^* dans $\vec{R}_B$ se conserve. On a donc deux intégrales du mouvement $\vec{\sigma}^*$ et E^* .

Le calcul de ces deux grandeurs constantes peut se faire à $t = 0$ ce qui donne :

$$\vec{\sigma}^* = \ell_1 \vec{u}_x \wedge \frac{m}{2} \ell_1 \omega \vec{u}_y = \frac{m}{2} \ell_1^2 \omega \vec{u}_z \text{ et } E^* = \frac{m}{4} \ell_1^2 \omega^2 + \frac{1}{2} k (\ell_1 - \ell_0)^2$$

Dans le cas $k = 2m\omega^2$ on trouve :

$$\vec{\sigma}^* = \frac{m}{2} (L + 2\ell_0)^2 \omega \vec{u}_z \text{ et } E^* = \frac{m}{4} (L + 2\ell_0)^2 \omega^2 + m\omega^2 (L + \ell_0)^2$$

- f) Le mouvement est plan, donc en coordonnées polaires la vitesse s'écrit

$$\vec{v} = \dot{r} \vec{u}_r + r \dot{\theta} \vec{u}_\theta$$

Le moment cinétique est donc :

$$\vec{\sigma}^* = \mu r^2 \dot{\theta} \vec{u}_z = \frac{m}{2} \ell_1^2 \omega \vec{u}_z \Rightarrow r^2 \dot{\theta} = \ell_1^2 \omega$$

L'énergie s'écrit donc :

$$E^* = \frac{1}{2} \mu (\dot{r}^2 + r^2 \dot{\theta}^2) + E_p = \frac{1}{2} \mu \dot{r}^2 + \frac{m \ell_1^4 \omega^2}{4 r^2} + \frac{k}{2} (r - \ell_0)^2 \text{ et } E_{p,\text{eff}} = \frac{m \ell_1^4 \omega^2}{4 r^2} + \frac{k}{2} (r - \ell_0)^2$$

La dérivée de l'énergie potentielle effective est :

$$\frac{dE_{p,\text{eff}}}{dr} = -\frac{m \ell_1^4 \omega^2}{2 r^3} + k(r - \ell_0)$$

Cette dérivée s'annule pour $k(r_m - \ell_0) = \frac{m \ell_1^4 \omega^2}{2 r_m^3}$ et r_m correspond à un minimum. L'expression de la dérivée de $\frac{dE_{p,\text{eff}}}{dr}(\ell_1)$ lorsque $k = 2m\omega^2$ est :

$$\frac{dE_{p,\text{eff}}}{dr}(\ell_1) = m\omega^2 \left(\frac{3}{2} L + \ell_0 \right) > 0$$

Comme r_m correspond à un minimum et que $\frac{dE_{p,\text{eff}}}{dr}(\ell_1) > 0$ cela signifie que $\ell_1 > r_m$.

L'allure de la fonction $E_{p,\text{eff}}$ est tracée sur la figure (S29.2).

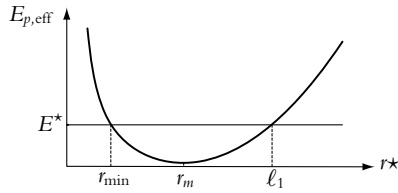


Figure S29.2

- g) À $t = 0$, le ressort a une longueur ℓ_1 et $\dot{r}^* = 0$, ce qui permet de positionner la droite horizontale représentant l'énergie mécanique E^* , elle croise la courbe en $r = \ell_1$. Le mouvement peut se faire dans la zone $E^* \geq E_{p,\text{eff}}$ puisque $m\dot{r}^2/2 \geq 0$, la valeur maximale de r^* est ℓ_1 d'après la figure.
- h) Dans $\mathcal{R}_B$, on assiste à une rotation autour de G avec le ressort qui oscille entre une longueur minimale $r_{\min}$ et une longueur maximale ℓ_1 . Dans $\mathcal{R}$, il faut composer ce mouvement avec le mouvement rectiligne de G .

Chapitre 30

A.1 1. Le système constitué du bateau et des passagers est soumis au poids et à la réaction de l'eau dans

le référentiel terrestre galiléen. Le théorème de la résultante dynamique s'écrit $M \frac{d\vec{v}_G}{dt} = \sum \vec{f}$. La projection sur l'horizontale de la résultante des forces est nulle : on a donc conservation de la quantité de mouvement du système.

Pour $t < 0$, la quantité de mouvement est $(2m_p + m_b) \vec{v}_0$ en notant m_p la masse d'un passager et m_b celle du bateau. Pour $t > 0$, elle vaut $(m_p + m_b) \vec{v}$. La conservation de la quantité de mouvement donne $\vec{v} = \frac{2m_p + m_b}{m_p + m_b} \vec{v}_0$ dont l'application numérique donne $v = 2,74 \text{ m.s}^{-1}$.

2. Dans ce cas, la conservation de la quantité de mouvement s'écrit

$$(2m_p + m_b) \vec{v}_0 = (m_p + m_b) \vec{v} + m_p \vec{v}_1 \quad \text{ce qui donne} \quad \vec{v} = \frac{2m_p + m_b}{m_p + m_b} \vec{v}_0 - \frac{m_p}{m_p + m_b} \vec{v}_1$$

Les vitesses $\vec{v}_1$ et $\vec{v}_0$ sont de sens opposé, $\|\vec{v}\| = 3,84 \text{ m.s}^{-1}$.

3. On a toujours $(2m_p + m_b) \vec{v}_0 = (m_p + m_b) \vec{v} + m_p \vec{v}_1$ mais les vitesses $\vec{v}_1$ et $\vec{v}_0$ sont perpendiculaires. En utilisant le théorème de Pythagore, on obtient

$$\|\vec{v}\| = \sqrt{\left(\frac{(2m_p + m_b)v_0}{m_p + m_b}\right)^2 + \left(\frac{m_p v_1}{m_p + m_b}\right)^2} = 2,95 \text{ m.s}^{-1}$$

et l'angle α avec $\vec{v}_0$ défini par

$$\tan \alpha = \frac{m_p v_1}{(2m_p + m_b) v_0} = 0,40 \quad \text{soit} \quad \alpha = 22^\circ.$$

A.2 Le système étant isolé, le centre de gravité de l'ensemble est immobile et le référentiel galiléen considéré est le référentiel barycentrique. On note M_1 et M_2 les deux étoiles.

La relation de définition du centre de gravité donne :

$$m_1 \overrightarrow{GM_1} + m_2 \overrightarrow{GM_2} = \vec{0} \Rightarrow x = \frac{r_1}{r_2} = \frac{m_2}{m_1} = 2$$

Si l'on définit la particule fictive M par $\overrightarrow{GM} = \overrightarrow{M_1M_2}$, la force exercée sur M est celle exercée par M_1 sur M_2 donc $\vec{f} = -Gm_1m_2 \frac{\vec{u}_r}{r^2}$ avec $r = GM$ et $\vec{u}_r = \frac{\overrightarrow{GM}}{r}$. La relation fondamentale de la dynamique s'écrit :

$$-\mu L \dot{\theta}^2 = -\frac{Gm_1m_2}{L^2} \Rightarrow (m_1 + m_2) = \frac{L^3 \dot{\theta}^2}{G} = \frac{L^3 4\pi^2}{GT^2}$$

où $\mu = m_1m_2/(m_1 + m_2)$ est la masse réduite. Des deux relations entre m_1 et m_2 , on déduit :

$$m_1 = \frac{L^3 4\pi^2}{3GT^2} = 1,23.10^{33} \text{ kg et } m_2 = \frac{L^3 8\pi^2}{3GT^2} = 2,46.10^{33} \text{ kg}$$

A.3 1. Le système est isolé et initialement immobile donc le centre de gravité est immobile. Le référentiel barycentrique $\mathcal{R}^*$ est immobile par rapport au référentiel galiléen de référence et il est lui-même galiléen. On définit la particule fictive M dans $\mathcal{R}^*$ par $\overrightarrow{GM} = \overrightarrow{AB}$.

La force ressentie par M (celle exercée par A sur B) est colinéaire à $\overrightarrow{GM}$, donc son moment par rapport à G est nul. Le théorème du moment cinétique entraîne :

$$\frac{d\vec{L}^*}{dt} = \vec{0} \Rightarrow \vec{L}^* = c\vec{e}$$

Initialement, la particule fictive est immobile comme A et B donc son moment cinétique par rapport à G est nul, on en déduit donc que $\vec{L}^* = \vec{0}$ pour tout t : le mouvement est rectiligne, le long de la droite AB .

Le force qui s'exerce sur M est une force électrostatique conservative, donc l'énergie mécanique E_m^* se conserve. Elle est égale à :

$$E_m^* = \frac{1}{2}\mu v^{*2} + \frac{q^2}{4\pi\epsilon_0 r}$$

On écrit l'égalité de cette énergie à l'instant initial et lorsque les particules sont infiniment éloignées :

$$\frac{q^2}{4\pi\epsilon_0 D} \frac{1}{D} = \frac{1}{2}\mu v_\infty^2 \Rightarrow v_\infty = \sqrt{\frac{q^2}{2\pi\epsilon_0 D\mu}} = \sqrt{\frac{q^2}{\pi\epsilon_0 D}}$$

où l'on a remplacé la masse fictive μ par sa valeur $1/2$.

Il reste à calculer les vitesses finales de A et B . Si l'on oriente la droite AB positivement dans le sens AB (vecteur unitaire $\vec{u}$), la vitesse finale de M est $\vec{v}^* = v_\infty \vec{u}$. Les expressions établies dans la réduction du problème à deux corps sont :

$$\vec{v}_A = -\frac{m_B}{m_A + m_B} \vec{v}^* \text{ et } \vec{v}_B = \frac{m_A}{m_A + m_B} \vec{v}^*$$

On en déduit que les vitesses finales des particules sont :

$$\vec{v}_B = -\vec{v}_A = \frac{1}{2} \sqrt{\frac{q^2}{\pi\epsilon_0 D}} \vec{u}$$

2. Dans ce deuxième cas, le système étant isolé mais la vitesse initiale de G n'étant pas nulle, le référentiel barycentrique $\mathcal{R}^*$ est en translation rectiligne uniforme par rapport au référentiel galiléen de référence. Le référentiel $\mathcal{R}^*$ est donc aussi galiléen. La vitesse est :

$$\vec{v}_G = \frac{m_B \vec{V}_0}{m_A + m_B} = \frac{1}{2} \vec{V}_0$$

On se place maintenant dans le référentiel $\mathcal{R}^*$.

Pour la même raison que dans le cas précédent, le mouvement a lieu sur la droite AB et l'on prend les mêmes notations.

La vitesse initiale de la particule fictive est :

$$\vec{v}_i^* = \vec{v}_B(0) - \vec{v}_A(0) = -V_0 \vec{u}$$

La distance minimale entre les deux charges correspond à la distance minimale de la particule fictive par rapport à G . En ce point, la vitesse de M va s'annuler et M va repartir dans l'autre sens en s'éloignant de G . On écrit comme précédemment, l'égalité des énergies mécaniques à l'instant initial

et à la position la plus proche de G :

$$\frac{1}{2}\mu V_0^2 + \frac{q^2}{4\pi\epsilon_0} \frac{1}{D} = \frac{q^2}{4\pi\epsilon_0} \frac{1}{d_{\min}} \Rightarrow d_{\min} = \left(\frac{1}{D} + \frac{V_0^2 \pi \epsilon_0}{q^2} \right)^{-1}$$

A.4 On appelle $\mathcal{R}$ le référentiel galiléen lié au plan xOy .

1. Dans le cas présent, le système (A, B, tige) n'est pas isolé mais pseudo-isolé. Les forces extérieures qui s'exercent sur lui se compensent. Il s'agit du poids de B et des réactions normales du plan. La résultante des forces extérieures est nulle donc :

$$\frac{d\vec{v}_{G/\mathcal{R}}}{dt} = \vec{0} \Rightarrow \vec{v}_{G/\mathcal{R}} = \vec{c\epsilon}$$

La vitesse initiale de G est :

$$\vec{v}_G = \frac{m_A \vec{V}_A(0)}{m_A + m_B} = \frac{v_0}{2} \vec{u}_y$$

Le mouvement de G est donc rectiligne uniforme.

2. On définit le mobile fictif M par $\vec{AB} = \vec{GM}$. La barre est rigide donc la distance $AB = GM$ est une constante. La force exercée sur M est la tension de la barre exercée sur B ; elle est donc colinéaire à $\vec{GM}$. On en déduit que le moment de cette force en G est nul et que le moment cinétique $\vec{L}_B$ dans $\mathcal{R}_B$ est une constante. On calcule $\vec{L}_B$ à $t = 0$. La vitesse initiale de M est égale à $\vec{v}_B(0) - \vec{v}_A(0)$ soit $-\nu_0 \vec{u}_y$ et à $t = 0$, $\vec{GM} = L \vec{u}_x$. On en déduit le moment cinétique :

$$\vec{L}_B = \vec{GM} \wedge \mu \vec{v}_{G/\mathcal{R}_B} = -\frac{1}{2} \nu_0 L u_z \vec{e}$$

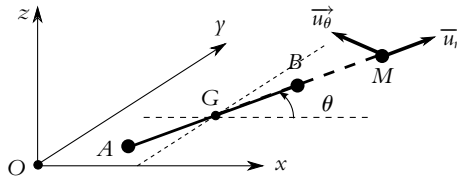


Figure S30.1

Sur la figure (S30.1), on a représenté la trajectoire rectiligne de G parallèle à Oy , ainsi que la position de la barre à un instant t . On repère le point M dans le référentiel barycentrique par des coordonnées polaires de repère $(\vec{u}_r, \vec{u}_\theta)$. Le point M décrit donc un mouvement circulaire autour de G à la vitesse angulaire ω telle que :

$$\vec{L}_B = L \vec{u} \wedge m L \omega \vec{u}_\theta = m L^2 \omega \vec{u}_z \Rightarrow \omega = -\frac{V_0}{2L}$$

Donc la tige tourne dans $\mathcal{R}_B$ dans le sens horaire autour du point G.

B.1 1. La force d'interaction s'exerçant entre le noyau et le proton est la force de Coulomb :

$\vec{f} = -\frac{e^2}{4\pi\epsilon_0 (NE)^3} \vec{NE}$ en notant N la position du noyau et E celle de l'électron. On peut négliger la force d'origine gravitationnelle; en comparant leur intensité respective, on a :

$$\frac{e^2}{4\pi\epsilon_0 G m_e m_p} = 2,3 \cdot 10^{29} \gg 1.$$

2. En projetant le principe fondamental de la dynamique sur la droite NE (dirigée par le vecteur $\vec{u}_r$ des coordonnées polaires), on obtient $m_e \omega^2 r = \frac{e^2}{4\pi\epsilon_0 r^2}$ avec $v = \omega r$ donc : $v^2 = \frac{e^2}{4\pi\epsilon_0 m_e r}$.

3. Le module du moment cinétique s'écrit : $L = m_e r v = \frac{nh}{2\pi}$. Par identification des deux expressions donnant le carré de la vitesse, on obtient : $v^2 = \frac{e^2}{4\pi\epsilon_0 m_e r} = \frac{n^2 h^2}{4\pi^2 m_e^2 r^2}$ soit $r = \frac{n^2 h^2 \epsilon_0}{\pi m_e e^2}$.

4. Application numérique pour $n = 1$: $r = 53$ pm

5. L'énergie mécanique vaut : $E_m = E_c + E_p = \frac{1}{2} m_e v^2 - \frac{e^2}{4\pi\epsilon_0 r} = -\frac{e^2}{8\pi\epsilon_0 r}$. De plus, $r = \frac{n^2 \hbar^2 \epsilon_0}{\pi m_e e^2}$ donc

$$E = -\frac{e^4 m_e}{8\epsilon_0^2 \hbar^2} \frac{1}{n^2} = -\frac{A}{n^2} \text{ avec } A = \frac{e^4 m_e}{8\epsilon_0^2 \hbar^2} = 2,17 \cdot 10^{-18} \text{ J} = 13,6 \text{ eV.}$$

6. $\Delta E = A \left(\frac{1}{n_2^2} - \frac{1}{n_1^2} \right) = h\nu = h \frac{c}{\lambda}$ donc $\frac{1}{\lambda} = \frac{A}{hc} \left(\frac{1}{n_2^2} - \frac{1}{n_1^2} \right) = R_H \left(\frac{1}{n_2^2} - \frac{1}{n_1^2} \right)$ avec $R_H = \frac{A}{hc} = \frac{e^4 m_e}{8\epsilon_0^2 hc^3} = 1,09 \cdot 10^7 \text{ m}^{-1}$.

7. Les longueurs d'onde visibles sont comprises entre 400 et 800 nm environ.

Pour la série de Lyman ($n_2 = 1$), on a : pour $n_1 = 2$, $\lambda = 122$ nm ; pour $n_1 = 3$, $\lambda = 103$ nm et pour toutes les valeurs supérieures de n_1 les longueurs d'onde seront plus faibles. On n'a donc aucune raie dans le domaine visible, elles sont toutes dans le domaine ultra-violet.

Pour la série de Balmer ($n_2 = 2$), on a : pour $n_1 = 3$, $\lambda = 660$ nm ; pour $n_1 = 4$, $\lambda = 489$ nm ; pour $n_1 = 5$, $\lambda = 436$ nm ; pour $n_1 = 6$, $\lambda = 412$ nm ; pour $n_1 = 7$, $\lambda = 399$ nm et pour toutes les valeurs supérieures de n_1 les longueurs d'onde seront plus faibles. On a donc quatre raies (pour n_1 égal à 3, 4, 5 et 6) dans le domaine visible.

B.2 1. Généralités :

- a) Il existe des référentiels particuliers appelés galiléens dans lesquels un point matériel isolé (non soumis à des interactions) a un mouvement rectiligne uniforme.
- b) Le référentiel barycentrique est le référentiel en translation à la vitesse du barycentre par rapport au référentiel galiléen de référence.
- c) L'étoile double est un système isolé donc en appliquant le principe fondamental de la dynamique, on établit que la vitesse du centre d'inertie G est constante et le référentiel barycentrique est galiléen.

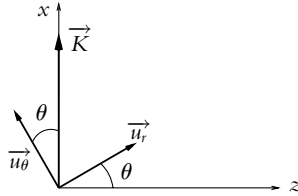
2. Mouvement des deux éléments dans le référentiel barycentrique :

- a) Dans le référentiel barycentrique, on a $m_1 \vec{v}_1^* + m_2 \vec{v}_2^* = \vec{0}$ et la vitesse relative s'écrit $\vec{v} = \vec{v}_2^* - \vec{v}_1^*$. On en déduit en résolvant le système en $\vec{v}_1^*$ et $\vec{v}_2^*$: $\vec{p}_1^* = -\vec{p}_2^* = \mu \vec{v}$ où $\mu = \frac{m_1 m_2}{m_1 + m_2}$ est la masse réduite du système.
- b) Par définition, $\vec{L}_G^* = \overrightarrow{GM_1} \wedge m_1 \vec{v}_1^* + \overrightarrow{GM_2} \wedge m_2 \vec{v}_2^*$. D'après la question précédente, $\vec{L}_G^* = (\overrightarrow{GM_2} - \overrightarrow{GM_1}) \wedge \mu \vec{v} = \overrightarrow{M_1 M_2} \wedge \mu \vec{v} = \overrightarrow{GM} \wedge \mu \vec{v}$ en posant $\overrightarrow{GM} = \overrightarrow{M_1 M_2}$.
- c) $E_c^* = \frac{1}{2} m_1 v_1^{*2} + \frac{1}{2} m_2 v_2^{*2} = \frac{1}{2} \mu v^2$.
- d) On écrit le principe fondamental de la dynamique pour chacun des constituants de l'étoile double dans le référentiel barycentrique. On en déduit : $\mu \frac{d\vec{v}}{dt} = \vec{f}_{1 \rightarrow 2}$.
- e) Par définition du centre d'inertie G , $m_1 \overrightarrow{GM_1} + m_2 \overrightarrow{GM_2} = \vec{0}$ donc $\overrightarrow{GM_1} = -\frac{m_2}{m_1 + m_2} \overrightarrow{GM}$ et $\overrightarrow{GM_2} = \frac{m_1}{m_1 + m_2} \overrightarrow{GM}$. Par conséquent, les trajectoires de M_1 et M_2 s'obtiennent à partir de celle de M par homothétie.

3. Étude du mouvement de M :

- a) $\frac{d\vec{L}^*}{dt} = \frac{d\overrightarrow{GM}}{dt} \wedge \mu \vec{v} + \overrightarrow{GM} \wedge \mu \frac{d\vec{v}}{dt}$. Les deux termes sont nuls car les produits vectoriels concernent des vecteurs colinéaires puisque $\frac{d\overrightarrow{GM}}{dt} = \vec{v}$ et $\mu \frac{d\vec{v}}{dt} = \vec{f}_{1 \rightarrow 2}$ est parallèle à $\overrightarrow{GM}$. Le moment cinétique est donc une constante : s'il est nul, le mouvement est rectiligne *a fortiori* plan et sinon il s'effectue dans le plan défini par G , M_0 et $\vec{v}_0$.

- b) On exprime le moment cinétique dans les coordonnées polaires du plan dans lequel a lieu le mouvement : $L = \mu r^2 \dot{\theta}$. L est une constante d'après ce qui précède donc la vitesse aréolaire $C = \frac{dA}{dt} = r^2 \dot{\theta}$ est constante. La loi des aires est vérifiée : l'aire balayée par le rayon vecteur est constante au cours du temps.
- c) Comme $\vec{u}_r = -\frac{1}{\dot{\theta}} \frac{d\vec{u}_\theta}{dt}$ et $\vec{f}_{1 \rightarrow 2} = -G \frac{m_1 m_2}{r^2} \vec{u}_r$, on en déduit le résultat en reportant dans l'expression du principe fondamental de la dynamique.
- d) On intègre cette relation $\vec{v} = \frac{G m_1 m_2}{C \mu} \vec{u}_\theta + \vec{K}$ en notant $\vec{K}$ la constante d'intégration.
- e) On multiplie scalairement cette relation par $\vec{u}_\theta$: $\vec{v} \cdot \vec{u}_\theta = \frac{G m_1 m_2}{C \mu} + \vec{K} \cdot \vec{u}_\theta$.



En choisissant l'axe Gx suivant $\vec{K}$, on obtient $\vec{K} \cdot \vec{u}_\theta = K \cos \theta$ et $r = \frac{p}{1 + e \cos \theta}$ avec $e = \frac{\mu C K}{G m_1 m_2}$.

- f) D'après la question précédente, on a $p = \frac{\mu C^2}{G m_1 m_2} = \frac{C^2}{G m}$.
- g) $r_1 = -\frac{m_2}{m_1 + m_2} \frac{p}{1 + e \cos \theta}$ et $r_2 = \frac{m_1}{m_1 + m_2} \frac{p}{1 + e \cos \theta}$ avec $\theta_1 = \theta_2 = \theta$.
On déduit des questions précédentes :

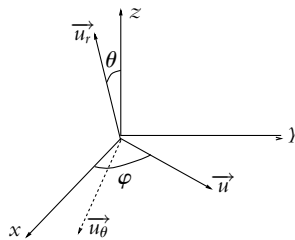
$$\vec{v}_1 = \dot{r}_1 \vec{u}_r + r_1 \dot{\theta} \vec{u}_\theta = -\frac{m_2}{m_1 + m_2} \frac{p \dot{\theta}}{1 + e \cos \theta} \left(\frac{e \sin \theta}{1 + e \cos \theta} \vec{u}_r + \vec{u}_\theta \right)$$

et

$$\vec{v}_2 = \frac{m_1}{m_1 + m_2} \frac{p \dot{\theta}}{1 + e \cos \theta} \left(\frac{e \sin \theta}{1 + e \cos \theta} \vec{u}_r + \vec{u}_\theta \right)$$

4. Observation du phénomène :

- a) $\vec{u} = \cos \varphi \vec{u}_x + \sin \varphi \vec{u}_y = \cos \varphi \sin \theta \vec{u}_r + \cos \varphi \cos \theta \vec{u}_\theta + \sin \varphi \vec{u}_y$.



- b) On explicite les expressions compte tenu des résultats de la question précédente.

$$\lambda_1 = \lambda_0 \left(1 + \frac{m_2 C^2 \dot{\theta} \cos \varphi (e + \cos \theta)}{G c (m_1 + m_2)^2 (1 + e \cos \theta)^2} \right) \quad \text{et} \quad \lambda_2 = \lambda_0 \left(1 - \frac{m_1 C^2 \dot{\theta} \cos \varphi (e + \cos \theta)}{G c (m_1 + m_2)^2 (1 + e \cos \theta)^2} \right)$$

- c) Dans l'expression obtenue, seul θ dépend du temps t donc il faut nécessairement $e = 0$ pour avoir une variation sinusoïdale. r est alors constant et la trajectoire est circulaire. En utilisant la loi des aires, on en déduit que $\omega = \dot{\theta}$ est une constante.

- d) $\epsilon_1 = \frac{m_2 C^2 \omega \cos \varphi}{Gc(m_1 + m_2)^2}$ et $\epsilon_2 = \frac{m_1 C^2 \omega \cos \varphi}{Gc(m_1 + m_2)^2}$.
- e) Si $\epsilon_1 = \epsilon_2$, $m_1 = m_2 = \frac{m}{2}$.
- f) On remplace ω par $\frac{2\pi}{T}$ et C par $r^2 \omega$ avec $r = \frac{\nu T}{2\pi}$. On obtient $\vec{v}_1 = -\sqrt[3]{\frac{4\pi Gm}{T}} \vec{u}_\theta = -\vec{v}_2$.
- g) On fait de même dans l'expression de ϵ et on tire $m = \frac{T}{16\pi^4 G} \left(\frac{\epsilon c}{\cos \varphi} \right)^3$.
- h) On obtient $m = 2,7 \cdot 10^{30}$ kg de l'ordre de la masse solaire et $v_1 = v_2 = 54,1 \text{ km.s}^{-1}$.
- i) $\rho = M_1 M_2 = \frac{\nu_1 T}{\pi} = \sqrt[3]{\frac{4GmT^2}{\pi^2}}$.
- j) L'application numérique donne $\rho = 245 \cdot 10^9 \text{ m}$.

B.3 1. La force d'interaction dérive de l'énergie potentielle. Si l'on appelle $\vec{u}_r$ le vecteur unitaire $\frac{\vec{r}}{r}$, on a :

$$\vec{F} = \vec{F}_{1 \rightarrow 2} = -\vec{F}_{2 \rightarrow 1} = -\frac{dE_p}{dr} \vec{u}_r = \left(\frac{12a}{r^{13}} - \frac{6b}{r^7} \right) \vec{u}_r$$

La position d'équilibre correspond à $\vec{F} = \vec{0}$ soit $r_e^6 = 2a/b = 2\sigma^6$ donc $r_e = 2^{1/6} \sigma$. On peut mettre la force sous la forme :

$$\vec{F} = \frac{6b}{r^7} \left(\frac{r_e^6}{r^6} - 1 \right) \vec{u}_r$$

Donc si $r < r_e$ la force est répulsive et si $r > r_e$ elle est attractive.

2. a) Le référentiel barycentrique est le référentiel lié au centre des gravité des particules et en translation par rapport au référentiel galiléen de référence $\mathcal{R}$.
- b) Dans le cas considéré, le système des deux particules est supposé isolé, donc la résultante des forces extérieures $\vec{F}_{\text{ext}}$ est nulle. Dans ce cas, on a :

$$(m_1 + m_2) \frac{d\vec{v}_{G/\mathcal{R}}}{dt} = \vec{F}_{\text{ext}} = \vec{0} \Rightarrow \vec{v}_{G/\mathcal{R}} = \vec{a}e$$

Le point G est en translation rectiligne uniforme dans $\mathcal{R}$ donc $\mathcal{R}_B$ est galiléen.

3. a) On utilise les deux relations suivantes :

$$m_1 \vec{r}_1 + m_2 \vec{r}_2 = \vec{0} \text{ et } \vec{r} = \vec{r}_2 - \vec{r}_1$$

ce qui entraîne :

$$\vec{r}_2 = \frac{m_1}{m_1 + m_2} \vec{r} \text{ et } \vec{r}_1 = -\frac{m_2}{m_1 + m_2} \vec{r}$$

- b) On passe d'une trajectoire à l'autre par une homothétie de centre G . Les trajectoires ont donc la même forme.
4. La quantité de mouvement de M_2 dans $\mathcal{R}_B$ est $\vec{p}_{2B} = m_2 \vec{v}_{2B}$, or la vitesse dans ce référentiel est la dérivée de $\vec{GM}_2 = \vec{r}_2$, donc :

$$\vec{p}_{2B} = \frac{m_1 m_2}{m_1 + m_2} \frac{d\vec{r}}{dt} = \mu \frac{d\vec{r}}{dt} = \vec{p}_M$$

5. a) On calcule la dérivée de $\vec{p}_M$:

$$\frac{d\vec{p}_M}{dt} = \mu \frac{d\vec{v}_M}{dt} = \frac{m_1 m_2}{m_1 + m_2} \left(\frac{d\vec{v}_2}{dt} - \frac{d\vec{v}_1}{dt} \right) = \frac{m_1 m_2}{m_1 + m_2} \left(\frac{\vec{F}_{1 \rightarrow 2}}{m_1} - \frac{\vec{F}_{2 \rightarrow 1}}{m_2} \right)$$

or $\vec{F}_{2 \rightarrow 1} = -\vec{F}_{1 \rightarrow 2}$, soit

$$\frac{d\vec{p}_M}{dt} = \mu \left(\frac{1}{m_2} + \frac{1}{m_1} \right) \vec{F}_{1 \rightarrow 2} = \mu \frac{1}{\mu} \vec{F}_{1 \rightarrow 2} = \vec{F}_{1 \rightarrow 2} = \vec{F}$$

- b) On applique le théorème du moment cinétique à M dans $\mathcal{R}_B$:

$$\frac{d\vec{\sigma}_b}{dt} = \vec{GM} \wedge \vec{F} = \vec{0}$$

car $\vec{F}$ et $\vec{GM}$ sont colinéaires. La dérivée du moment cinétique est nulle donc $\vec{\sigma}_b$ est un vecteur constant. Par définition du moment cinétique, il est perpendiculaire à $\vec{v}_M$ et $\vec{GM}$; ces deux vecteurs sont donc perpendiculaires à une direction constante pour tout t , donc ils sont toujours dans le même plan.

- c) L'expression de la vitesse en coordonnées polaires est $\vec{v} = \dot{r}\vec{u}_r + r\dot{\theta}\vec{u}_\theta$.

6. a) On calcule l'énergie mécanique de M égale à l'énergie mécanique du système (M_1, M_2) dans $\mathcal{R}_B$:

$$E_m = \frac{1}{2}\mu v^2 + E_p = \frac{1}{2}\mu(\dot{r}^2 + r^2\dot{\theta}^2) + E_p$$

or en coordonnées polaires le moment cinétique est :

$$\vec{\sigma}_b = \vec{GM} \wedge \mu \vec{v} = \mu r^2 \dot{\theta} \vec{u}_z$$

on peut remplacer $\dot{\theta}$ dans l'expression de l'énergie :

$$E_m = \frac{1}{2}\mu\dot{r}^2 + \frac{\sigma_b^2}{2\mu r^2} + E_p = \frac{1}{2}\mu\dot{r}^2 + E_{p,\text{eff}} \text{ avec } E_{p,\text{eff}} = \frac{\sigma_b^2}{2\mu r^2} + \frac{a}{r^{12}} - \frac{b}{r^6}$$

L'énergie mécanique peut donc s'exprimer comme la somme de l'énergie cinétique radiale et d'un terme ne dépendant que de r et assimilable à une énergie potentielle (dite effective).

- b) En remplaçant dans l'expression de E_p , a par $\epsilon\sigma^{12}$ et b par $\epsilon\sigma^6$, on obtient $r = \sigma x$.

- c) Sachant que σ est constant, on obtient $\dot{r} = \sigma\dot{x}$.

- d) La mise en forme de l'équation revient à poser $p = \frac{\sigma_b^2}{2\mu\epsilon\sigma^2}$.

- e) Si l'on remplace μ par $m/2$, on obtient : $\sigma_b = \sigma\sqrt{p\epsilon m}$. L'application numérique donne :

- $\sigma_b = 5,37 \cdot 10^{-33} \text{ kg.m}^2.\text{s}^{-1}$,
- $\sigma_b = 1,07 \cdot 10^{-32} \text{ kg.m}^2.\text{s}^{-1}$.

7. Dans toute la discussion, on utilise le fait que la zone de mouvement est telle que $E_m \geq E_{p,\text{eff}}$ (puisque $m\dot{r}^2/2 \geq 0$).

- a) Pour $E_m = E_0$, la seule possibilité est l'égalité $E_m = E_{p,\text{eff}}$. Donc $x = x_0$ et $r = r_0 = \sigma x_0$. De plus, $\dot{x} = 0$ et $\dot{r} = 0$. Le point M a donc un mouvement circulaire de centre G et uniforme car $\dot{\theta} = C/(\mu r_0^2)$ est constant.

- b) Sur la figure (S30.2) on a agrandi la partie négative de la courbe. On a tracé la droite horizontale représentant l'énergie mécanique E_1 .

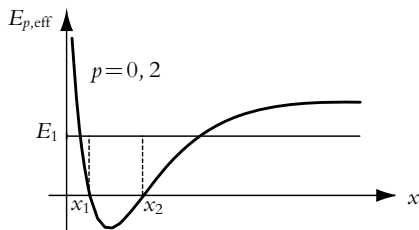


Figure S30.2

La partie de la courbe qui correspond à $E_1 \geq E_{p,\text{eff}}$ correspond à une zone de l'espace pour laquelle $x_1 \leq x \leq x_2$. La particule M est piégée dans le puits de potentiel et va osciller entre ces deux positions extrêmes. C'est ce que l'on constate aussi sur les trajectoires de phase correspondantes (par exemple celles correspondant à $E_m = -0,05$ ou $E_m = 0$) qui sont des trajectoires de phases fermées. On a alors un état lié.

- c) Sur la figure (S30.3) on a agrandi la partie de la courbe correspondant $0 \leq x \leq 4$. On a tracé la droite horizontale représentant l'énergie mécanique E_2 .

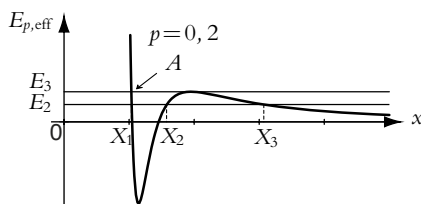


Figure S30.3

La partie de la courbe qui correspond à $E_1 \geq E_{p,eff}$ correspond à une zone de l'espace pour laquelle $X_1 \leq x \leq X_2$ ou $x \geq X_3$. Le mouvement va donc dépendre de la position initiale, car la barrière de potentiel (hauteur E_3) n'est pas franchissable puisque $E_2 < E_3$.

- Si initialement, la particule M se trouve à une position $x < x_3$, c'est-à-dire à gauche de la barrière de potentiel, alors elle est piégée dans le puits de potentiel, comme dans la question précédente et elle se déplace entre X_1 et X_2 : on a alors un état lié.
- Si initialement, la particule M se trouve à une position $x > x_3$, c'est-à-dire à droite de la barrière de potentiel, alors elle ne pourra s'approcher qu'à une valeur de x égale à X_3 , et comme elle dans la zone où la force effective est répulsive (courbe de $E_{p,eff}$ décroissante) elle va s'éloigner de G jusqu'à l'infini où son énergie cinétique radiale $mr^2/2$ sera égale à E_2 puisque $E_{p,eff} \rightarrow \infty$: on a alors un état de diffusion.

- d) On a tracé sur la figure (S30.3), la droite horizontale représentant l'énergie mécanique E_3 . La position $x = x_3$ correspond à une position d'équilibre instable, donc lorsque M passe à la position $x = x_3$, il peut soit se rapprocher de G (x diminue) soit s'éloigner de G (x augmente). Le point le plus proche de G est le point d'intersection A (repéré par la flèche) avec la courbe $E_{p,eff}$.

Si M est lancé de $x > x_3$ vers G , il se rapproche de G . Premier cas, en $x = x_3$ il fait demi-tour (il atteint cette position avec une vitesse nulle puisque $E_m = E_{p,eff}$) et repart vers l'infini : il s'agit d'un état de diffusion. Deuxième cas, il passe la barrière de potentiel et atteint le point A où il rebrousse chemin, il franchit de nouveau la barrière de potentiel et repart vers l'infini : c'est encore un état de diffusion. Troisième cas, la position initiale du point est entre x_A et x_3 dans le puits de potentiel et M oscille de x_A à x_3 : c'est un état lié, mais physiquement il faut qu'il ait une énergie légèrement inférieure à E_3 sinon M peut franchir la barrière de potentiel. Enfin dernier état non réalisable physiquement, M est initialement en $x = x_3$ sans vitesse : c'est un état d'équilibre mais il est instable.

- e) Pour $E_m = E_4$, la barrière de potentiel en $x = x_3$ est franchissable, donc quelle que soit la position initiale de M , le point finira à l'infini, c'est un état de diffusion. Les trajectoires de phase sont les courbes non fermées comme celle $E_m = 0, 2$.
- f) On observe dans tous les cas un état de diffusion, la position de plus grande approche de G correspondant à l'intersection de la droite horizontale E_m avec la courbe $E_{p,eff}$. Les trajectoires de phases sont toutes des courbes ouvertes.

Chapitre 31

- A.1** 1. On utilise la troisième loi de Képler soit $a = \left(\frac{T^2 GM_s}{4\pi^2} \right)^{1/3}$. L'application numérique donne $a = 1,26 \cdot 10^{13}$ m. L'équation de la trajectoire donne la distance au périhélie $d = p/(1 + e)$. Avec la

relation entre p et a : $p = a(1 - e^2)$ on en déduit :

$$a = \frac{d}{1 - e} \Rightarrow e = 1 - \frac{d}{a} = 0,9999 \text{ et } p = d(1 + e) = 2,32 \cdot 10^9 \text{ m}$$

Pour calculer les vitesses, on utilise le fait que l'aphélie et le périhélie sont les deux seuls points de la trajectoire pour lesquels la vitesse est perpendiculaire au rayon vecteur donc si on nomme C la constante des aires $C = r_A v_A$ et $C = r_P v_P$.

Lors de la résolution de l'équation du mouvement, on établit que $p = C^2/(GM_s)$. On peut donc écrire :

$$v_A = \frac{C}{r_A} = \sqrt{pGM_s} \frac{1 - e}{p} = 24 \text{ m.s}^{-1} \text{ et } v_P = \frac{C}{r_P} = \sqrt{pGM_s} \frac{1 + e}{p} = 480 \text{ km.s}^{-1}$$

L'excentricité est très proche de 1, la trajectoire est donc presque parabolique. Une perturbation apportée à la trajectoire lorsque la comète passe près d'une planète peut la rendre parabolique ou même hyperbolique.

2. a) La conservation de la quantité de mouvement entraîne $m \vec{v}_P = m_1 \vec{v}_1 + m_2 \vec{v}_2$ ce qui donne en projection parallèlement et perpendiculairement à $\vec{v}_P$:

$$m_1 v_1 \sin \alpha_1 = m_2 v_2 \sin \alpha_2 \text{ et } m v_P = m_1 v_1 \cos \alpha_1 + m_2 v_2 \cos \alpha_2$$

L'énoncé fait l'hypothèse $v_1 \simeq v_2$ donc $\frac{m_1}{m_2} = \frac{\sin \alpha_2}{\sin \alpha_1}$ et si l'on note v la vitesse commune des fragments :

$$v = \frac{m v_P}{m_1 \cos \alpha_1 + m_2 \cos \alpha_2} = v_P \frac{m_1 + m_2}{m_1 \cos \alpha_1 + m_2 \cos \alpha_2} = v_P \frac{\frac{m_1}{m_2} + 1}{\cos \alpha_1 + \cos \alpha_2}$$

L'application numérique donne $m_1/m_2 = 0,446$ et $v = 1,36 v_P$.

- b) L'expression de l'énergie mécanique par unité de masse est :

$$E_m = \frac{1}{2} v^2 - \frac{GM_s}{d} = -1,15 \cdot 10^{21} \text{ J}$$

L'énergie est négative, la trajectoire est donc une ellipse.

A.2 1. Sur l'orbite circulaire dans le référentiel géocentrique, la relation fondamentale de la dynamique donne :

$$-m r_0 \dot{\theta}^2 = -m \frac{v_c^2}{r_0} = -\frac{GmM_T}{r_0^2} \Rightarrow v_c = \sqrt{\frac{GM_T}{r_0}}$$

2. a) On suppose que la nouvelle vitesse est communiquée très rapidement, donc on néglige le déplacement du point. De plus, la vitesse communiquée est perpendiculaire au rayon vecteur $\vec{r}_0$ donc le point A est soit le périhélie (ellipse, parabole, hyperbole), soit l'apogée (ellipse).

Deux possibilités :

- Si A est le périhélie alors $r_0 = \frac{p}{1 + e}$.
- Si A est l'apogée alors $r_0 = \frac{p}{1 - e}$.

- b) Dans le cours, il a été établi avec l'équation du mouvement que $p = mC^2/(GmM_T)$. On calcule la constante des aires en A en rappelant qu'elle est égale au moment cinétique divisée par la masse, soit $C = r_0 v_A$. On a alors :

$$p = \frac{(r_0 v_A)^2}{GM_T} = r_0 \left(\frac{v_A}{v_c} \right)^2$$

On en déduit :

- Si A est le périhélie alors $e = \left(\frac{v_A}{v_c} \right)^2 - 1$.
- Si A est l'apogée alors $e = 1 - \left(\frac{v_A}{v_c} \right)^2$.

c) On examine les différents cas :

- Si $v_A = v_c$, les deux cas donnent $e = 0$, on retrouve évidemment une trajectoire circulaire.
- Si $v_A = v_C/3$, le premier cas donne $e < 0$ ce qui n'est pas physique et le deuxième cas donne $e = 0,91$. Il s'agit donc d'une trajectoire elliptique d'apogée A .
- Si $v_A = 3v_c$, le deuxième cas donne $e < 0$ ce qui n'est pas physique et le premier cas donne $e = 8$ ce qui correspond à une trajectoire hyperbolique de périégée A .

A.3 1. Si on repère la position du satellite par rapport à la direction de l'apogée, l'équation de la trajectoire est $r = \frac{p}{1 - e \cos \theta}$ (il y a un signe "+" au dénominateur si l'axe d'origine est en direction du périégée). Le satellite retombe sur Terre après avoir tourné de 90° donc pour $\theta = 90^\circ$ on a $r = R_T$, d'où $p = R_T$.

D'autre part, comme l'origine de la courbe est A , alors $r_0 = p/(1 - e)$.

2. Il a été démontré dans le cours $E_m = -GM_T m/(2a)$. Des relations établies à la question précédente, on tire $e = 1 - \frac{R_T}{r_0}$ et $E_m = -\frac{GmM_T}{2} \frac{(2r_0 - R_T)}{r_0^2}$

3. Au cours du mouvement, l'énergie mécanique se conserve donc en A :

$$\frac{1}{2}mv_A^2 - \frac{GmM_T}{R_T} = E_m \Rightarrow v_A = \frac{\sqrt{GM_T R_T}}{r_0}$$

A.4 1. a) La force exercée par l'étoile sur la planète est :

$$\vec{f} = -\frac{GM_p M_e}{r^2} \vec{u}_r$$

b) Dans le cas d'un mouvement à force centrale le mouvement est plan (voir démonstration du cours). Le plan du mouvement est le plan perpendiculaire au moment cinétique et passant par le centre attracteur (ici le centre de l'étoile).

Par définition du moment cinétique :

$$\vec{L} = \vec{OP} \wedge M_p \vec{v}_p = M_p r \vec{u}_r \wedge (\dot{r} \vec{u}_r + r \dot{\theta} \vec{u}_\theta) = M_p r^2 \dot{\theta} \vec{u}_z$$

2. a) On calcule la dérivée du vecteur $\vec{e}$ sachant que L est constant :

$$\frac{d\vec{e}}{dt} = -\frac{L}{GM_e M_p} \vec{a} - \dot{\theta} \vec{u}_r = -\frac{M_p r^2 \dot{\theta}}{GM_e M_p} \left(-\frac{GM_e}{r^2} \vec{u}_r \right) - \dot{\theta} \vec{u}_r = \vec{0}$$

b) Étant donnée la direction de $\vec{e}$ le produit scalaire $\vec{e} \cdot \vec{u}_\theta$ est égal à $-e \cos \theta$. D'autre part, avec l'expression de $\vec{e}$, on obtient :

$$\vec{e} \cdot \vec{u}_\theta = -\frac{L}{GM_e M_p} \vec{v} \cdot \vec{u}_\theta + 1 = -\frac{L}{GM_e M_p} r \dot{\theta} + 1 = -\frac{L^2}{GM_e M_p^2} \frac{1}{r} + 1$$

En utilisant les deux expressions, on trouve :

$$-\frac{L^2}{GM_e M_p^2} \frac{1}{r} + 1 = -e \cos \theta \Rightarrow r = \frac{p}{1 + e \cos \theta} \text{ avec } p = \frac{L^2}{GM_e M_p^2}$$

c) Il apparaît dans l'expression précédente que la norme de e est égale à l'excentricité de la trajectoire. L'expression de p a été donnée à la question précédente.

d) Dans le cas d'un mouvement circulaire, l'excentricité est nulle.

e) Dans le cas d'un mouvement circulaire, la vitesse est perpendiculaire au rayon vecteur donc : $L = M_p R v_c$. Puisque dans ce cas $\vec{e} = \vec{0}$, on peut écrire :

$$\frac{L}{GM_p M_e} \vec{v}_c = \vec{u}_\theta \Rightarrow \frac{R}{GM_e} v_c \vec{v}_c = \vec{u}_\theta \Rightarrow \vec{v}_c = \sqrt{\frac{GM_e}{R}} \vec{u}_\theta$$

A.5 On étudie le système isolé constitué du point matériel P et de la planète dans le référentiel du centre de masse G supposé galiléen. Comme la masse de la planète M est très grande devant m , celle du point matériel, on peut confondre G avec O le centre de la planète. Dans ce référentiel, P est soumis à l'attraction gravitationnelle de la planète $\vec{F} = m\vec{G}$. Comme la planète a une répartition de masse à symétrie sphérique, $\vec{G}$ ne dépend que de r et est dirigé selon $\vec{u}_r$. En appliquant le théorème de Gauss, on a $4\pi r^2 G(r) = -4\pi \mathcal{G}M$ soit $G(r) = \frac{\mathcal{G}M}{r^2}$. En appliquant le principe fondamental de la dynamique au point P , on obtient $m \frac{d\vec{v}}{dt} = -\frac{\mathcal{G}Mm}{r^2} \vec{u}_r$. La seule force étant radiale, le mouvement a lieu selon $\vec{u}_r$. En projetant la relation précédente dans cette direction, on a $\ddot{r} = -\frac{\mathcal{G}M}{r^2}$. Par intégration (après avoir multiplié par $\dot{r}$), on en déduit $\frac{\dot{r}^2}{2} - 0 = -\mathcal{G}M \left(-\frac{1}{r} + \frac{1}{r_0} \right)$. En posant $v_0^2 = \frac{2\mathcal{G}M}{r_0}$ et $X = \frac{r}{r_0}$, on peut écrire $\dot{X}^2 = \frac{v_0^2}{r_0^2} \frac{1-X}{X}$. La durée cherchée s'obtient par intégration de cette relation soit $\frac{v_0}{r_0} \tau = - \int_1^{\frac{R}{r_0}} \sqrt{\frac{X}{1-X}} dX$. Cette intégrale peut être obtenue en posant $X = \cos^2 \varphi$. On en déduit $\tau = \frac{r_0}{v_0} \left(\text{Arccos} \sqrt{\frac{R}{r_0}} + \sqrt{\frac{R}{r_0}} \sqrt{1 - \frac{R}{r_0}} \right)$ soit pour $r_0 = 5R$ $\tau = \frac{5R}{v_0} \left(\text{Arccos} \sqrt{\frac{1}{5}} + \frac{2}{5} \right) = 11,9 \frac{R^{\frac{3}{2}}}{\sqrt{\mathcal{G}M}}$.

A.6 1. Dans le référentiel géocentrique, la seule force qui s'exerce est l'attraction gravitationnelle de la Terre qui est dirigée vers le centre T de la Terre. Elle est donc radiale et centrale de centre T et même centripète.

2. La force est centrale donc son moment par rapport à T est nul. Le théorème du moment cinétique implique donc que $\vec{L}_T(S)/\mathcal{R}$ est une constante vectorielle. Le mouvement a lieu dans un plan perpendiculaire à $\vec{L}_T(S)/\mathcal{R}$ et passant par T . Comme on ne connaît pas ce plan on ne connaît pas la direction de $\vec{L}_T(S)/\mathcal{R}$.

3. Par définition de l'accélération :

$$\vec{a}(S)/\mathcal{R} = \left[\frac{d\vec{v}(S)/\mathcal{R}}{dt} \right]_{\mathcal{R}}$$

L'expression de l'accélération en coordonnées polaires pour un mouvement circulaire de rayon $R_T + h$ est :

$$m \vec{a}(S)/\mathcal{R} = -m(R_T + h)\dot{\varphi}^2 \vec{u}_\rho = -m \frac{||\vec{v}(S)/\mathcal{R}||^2}{R_T + h} \vec{u}_\rho$$

4. La relation fondamentale de la dynamique appliquée au mouvement circulaire de rayon $R_T + h$ donne :

$$m \vec{a}(S)/\mathcal{R} = -\frac{mM_T \mathcal{G}}{(R_T + h)^2}$$

soit avec ce qui précède :

$$||\vec{v}(S)/\mathcal{R}|| = \sqrt{\frac{\mathcal{G}M_T}{R_T + h}}$$

donc aucune réponse n'est juste.

5. La période est T_0 est telle que :

$$T_0 = \frac{2\pi(R_T + h)}{||\vec{v}(S)/\mathcal{R}||} = 2\pi \frac{(R_T + h)^{3/2}}{\sqrt{\mathcal{G}M_T}}$$

Elle est donc indépendante de m .

6. L'application numérique donne pour $R_T + h$ environ 42 300 km soit pour h environ 35 900 km. La réponse à choisir est donc 36 000 km. Attention, il y a un piège dans l'énoncé puisque la masse de la Terre est donnée en grammes.

7. Avec l'expression de la vitesse, on calcule l'expression de l'énergie cinétique :

$$E_c = \frac{1}{2} m \| \vec{v}(S)_{/\mathcal{R}} \|^2 = \frac{m G M_T}{2(R_T + h)}$$

L'application numérique avec $h = 36\,000$ km donne $E_c = 9440$ MJ donc la réponse (B) est valable.

8. L'énergie potentielle :

$$E_p = -\frac{G M_T m}{R_T + h} \simeq -18,9 \text{ MJ}$$

donc la réponse (B) est valable.

9. On remarque que $E_p = -2E_c$. L'énergie mécanique $E_m = E_c + E_p$ est donc égale à $-E_c$ ce qui valide la réponse (C).

10. Par définition, le périégée est le point le plus proche du centre attractif. On sait de plus que l'énergie mécanique se conserve, or elle est égale à :

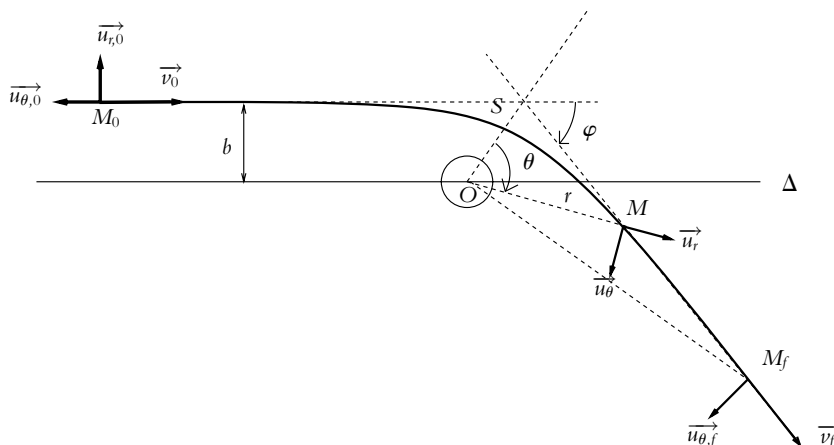
$$E_m = \frac{1}{2} m v^2 - \frac{G M_T m}{r} \Rightarrow v = \sqrt{\frac{2E_m}{m} + \frac{G M_T}{2r}}$$

donc lorsque r est minimal v est maximale.

11. La troisième loi de Képler donne :

$$\frac{T^2}{a^3} = \frac{4\pi^2}{G M_T} \Rightarrow a = \left(\frac{G M_T T^2}{4\pi^2} \right)^{1/3} = 155.10^3 \text{ km}$$

- B.1** 1. On étudie le météorite dans le référentiel géocentrique supposé galiléen. Sa trajectoire est une branche d'hyperbole de foyer O, le centre de la Terre.



La conservation du moment cinétique en O du météorite permet d'écrire, avec les notations de la figure ci-dessus : $\vec{L}_O = \vec{OM} \wedge m \vec{v} = -m r^2 \dot{\theta} \vec{u}_z = \text{constante}$. On calcule cette constante au "départ" du météorite : $\vec{L}_O = \vec{OM}_0 \wedge m \vec{v}_0 = -m b v_0 \vec{u}_z$ d'où $r^2 \dot{\theta} = b v_0$. Au sommet de la trajectoire (point S), $\vec{L}_O = -d v_S \vec{u}_z$, d'où

$$d v_S = b v_0 \quad (1)$$

La conservation de l'énergie entre les points M_0 et S permet d'obtenir une deuxième relation entre d et v_S :

$$\frac{1}{2} m v_0^2 = \frac{1}{2} m v_S^2 - \frac{G m M}{d} \quad (2)$$

En éliminant v_S entre les équations (1) et (2), on obtient l'équation vérifiée par d que l'on met sous la forme :

$$d^2 + 2 \frac{G M}{v_0^2} d - b^2 = 0$$

La résolution de cette équation donne :

$$d = -\frac{GM}{v_0^2} + \sqrt{\left(\frac{GM}{v_0^2}\right)^2 + b^2}$$

2. Le météorite ne rencontre pas la Terre si $d > R_T$ donc si $b > b_{\min}$ où $b_{\min} = R_T \sqrt{1 + 2\frac{GM}{R_T v_0^2}}$.
3. Le principe fondamental de la dynamique appliqué au météorite dans le référentiel géocentrique s'écrit : $m \frac{d\vec{v}}{dt} = -\frac{GmM}{r^2} \vec{u}_r$. Or $\frac{d\vec{u}_\theta}{dt} = -\dot{\theta} \vec{u}_r$ d'où

$$m \frac{d\vec{v}}{dt} = -\frac{GmM}{r^2 \dot{\theta}} \frac{d\vec{u}_\theta}{dt} = -\frac{GmM}{bv_0} \frac{d\vec{u}_\theta}{dt} \quad (3)$$

En intégrant cette équation entre le point de "départ" M_0 et un point loin à l'infini M_f , on obtient :

$$\vec{v}_0 - \frac{GM}{bv_0} \vec{u}_{\theta 0} = \vec{v}_f - \frac{GM}{bv_0} \vec{u}_{\theta f} \quad (4)$$

Les notations sont définies sur la figure ci-dessus. D'autre part, les points M_0 et M_f sont infiniment éloignés de la Terre, l'énergie potentielle d'interaction entre le météorite et la Terre y est nulle, l'énergie du météorite est, en ces points-là, uniquement sous forme cinétique. La conservation de l'énergie entre ces deux points donne : $v_0 = v_f$. On projette alors la relation (4) l'axe Ox :

$$v_0 = v_0 \cos \varphi + \frac{GM}{bv_0} \sin \varphi$$

On trouve tous calculs faits :

$$\tan \frac{\varphi}{2} = \frac{GM}{bv_0^2}$$

Si on avait projeté sur l'axe Oy , on aurait obtenu :

$$-\frac{GM}{bv_0} = -v_0 \sin \varphi + \frac{GM}{bv_0} \cos \varphi$$

Le lecteur est invité à vérifier qu'on obtient alors la même expression de $\tan \frac{\varphi}{2}$.

REMARQUE : En intégrant l'équation (3) entre le point M_0 et le sommet de la trajectoire, on obtient une relation entre v_0 et v_S qui, combinée à la conservation du moment cinétique, permet d'obtenir l'expression de d sans utiliser la conservation de l'énergie.

B.2 On étudie le satellite dans le référentiel géocentrique supposé galiléen. Dans ce référentiel, le satellite n'est soumis qu'à la force de gravitation colinéaire à son vecteur position.

1. La résultante des forces étant colinéaire avec le vecteur position, son moment par rapport à la Terre est nul et on en déduit par le théorème du moment cinétique que ce dernier est constant au cours du temps.
2. Si la valeur du moment cinétique est nulle, le mouvement est rectiligne. Sinon le mouvement a lieu dans le plan passant par O , centre de la Terre, et perpendiculaire au moment cinétique.
3. Dans ce plan, on utilise les coordonnées polaires et on explicite l'expression du moment cinétique soit $\vec{L}_0 = \vec{OM} \wedge m \vec{v} = r \vec{u}_r \wedge m (\dot{r} \vec{u}_r + r \dot{\theta} \vec{u}_\theta) = mr^2 \dot{\theta} \vec{u}_z$. On peut définir la constante des aires par $C = r^2 \dot{\theta}$.
4. Le mouvement est circulaire si r est une constante. Comme la constante C est une constante, on en déduit que si r est constant, $\dot{\theta}$ est également constant. Par conséquent, le mouvement est uniforme.
5. La vitesse s'exprime en coordonnées polaires par $\vec{v} = \dot{r} \vec{u}_r + r \dot{\theta} \vec{u}_\theta = r \dot{\theta} \vec{u}_\theta$ dans le cas d'un mouvement circulaire. La projection du principe fondamental de la dynamique sur $\vec{u}_r$ donne $-mr\dot{\theta}^2 = -\frac{GM_T m}{r^2}$ soit $C = r\dot{\theta} = \frac{GM_T}{r^2}$. On en déduit $v^2 = r^2 \dot{\theta}^2 = \frac{GM_T}{r}$.

6. L'énergie cinétique peut s'exprimer par $E_c = \frac{1}{2} m r^2 \dot{\theta}^2 = \frac{G m M_T}{2r}$.

7. Quant à l'énergie potentielle, elle vaut $E_p = -\frac{G m M_T}{r} = -2E_c$.

8. On en déduit l'énergie mécanique $E_m = E_c + E_p = -E_c = \frac{E_p}{2} = -\frac{G m M_T}{2r}$.

9. On effectue un développement limité de l'énergie potentielle

$$E_p(r + \Delta r) = -\frac{G m M_T}{r} \left(1 + \frac{\Delta r}{r}\right)^{-1} = -\frac{G m M_T}{r} \left(1 - \frac{\Delta r}{r}\right)$$

$$\text{On en déduit } \Delta E_p = \frac{G M_T m \Delta r}{r^2}.$$

10. Comme $E_m = \frac{E_p}{2}$ reste valable sur la trajectoire qui est supposée quasi-circulaire, on en déduit $\Delta E_m = \frac{G M_T m \Delta r}{2r^2}$.

11. Le travail de la force de frottement s'écrit $W_f = -\alpha m v \vec{v} \cdot \vec{v} T_0 = -\alpha m v^3 T_0$ car v est constante.

12. L'application du théorème de l'énergie mécanique donne $\frac{dE_m}{dt} = W_f$ soit en utilisant les résultats précédents et l'expression de la période de révolution $T_0 = \frac{2\pi r}{v}$, on obtient $\Delta r = -4\pi \alpha r^2$.

13. Le satellite se rapproche de la Terre sous le freinage de la force de frottement.

14. L'expression de la vitesse conduit à $v = \sqrt{\frac{G M_T}{r}} = \frac{2\pi r}{T_0}$. On en déduit la troisième loi de Kepler $\frac{T_0^2}{r^3} = \frac{4\pi^2}{G M_T}$.

15. Avec les hypothèses proposées dans l'énoncé, on a $T_0 dr = -4\pi \alpha r^2 dt$. En utilisant la troisième loi de Kepler, on a $dr = -2\alpha \sqrt{G M_T} r dt$. En intégrant entre r_0 et r pour une origine des temps en r_0 , on en déduit

$$\sqrt{r} = \sqrt{r_0} - \alpha \sqrt{G M_T} t \quad \text{soit} \quad K = -\alpha \sqrt{G M_T}.$$

B.3 1. a) Pour une trajectoire parabolique, l'équation en polaires donne $r_p = p/2$ où r_p est la position du périhélie. L'application numérique donne $p = 9,48 \cdot 10^{10}$ m en prenant 1 u.a. = $1,5 \cdot 10^{11}$ m.

b) On a démontré dans le cours la relation entre p et C : $p = C^2 / (M_s G)$. D'autre part, au périhélie $\vec{v}_p$ et $\vec{r}_p$ sont perpendiculaires donc $C = r_p v_p$. Après calcul, on trouve $v_p = \sqrt{\frac{2 G M_s}{r_p}}$. L'application numérique donne $v_p = 75,1 \text{ km.s}^{-1}$.

2. La constante des aires s'exprime en coordonnées polaires par $C = r^2 \dot{\theta}$ et en prenant l'origine des angles au périhélie $r = p / (1 + e \cos \theta)$ d'où :

$$\frac{d\theta}{dt} = \frac{C}{r^2} = \frac{C}{p^2} (1 + \cos \theta)^2$$

3. a) De l'équation précédente, on tire :

$$dt = \frac{p^2}{C} \frac{d\theta}{(1 + \cos \theta)^2} \Rightarrow \int_{t_0}^t dt = t - t_0 = \int_0^\alpha \frac{p^2}{C} \frac{d\theta}{(1 + \cos \theta)^2}$$

b) Le changement de variable entraîne $dx = (1 + \tan^2 \frac{\theta}{2}) d\theta$ soit $dx = (1 + x^2) d\theta$. D'autre part :

$$\frac{1}{(1 + \cos \theta)^2} = \frac{1}{4} \frac{1}{\cos^2 \frac{\theta}{2}} = \frac{1}{4} \left(1 + \tan^2 \frac{\theta}{2}\right)^2 = \frac{1}{4} (1 + x^2)^2$$

Finalement, l'intégrale devient :

$$t - t_0 = \frac{p^2}{2C} \int_0^X (1 + x^2) dx = \frac{p^2}{2C} \left(X + \frac{X^3}{3}\right)$$

- c) Il y a 151 jours entre le 8 novembre 1956 et le 8 avril 1957 (passage au périhélie). La résolution à la calculatrice de l'équation précédente donne $X = 2,826$ soit $\alpha = -140^\circ$ et $r = 4,25 \cdot 10^{11}$ m.

B.4 1. Si l'on néglige l'énergie gravitationnelle (potentielle) en A , l'énergie mécanique est égale à l'énergie cinétique soit $E_m = mv_A^2/2$. Puisque l'énergie mécanique est positive, la trajectoire est une branche enveloppante (force attractive) d'hyperbole de foyer C .

2. Dans le cours on a montré que pour une force centrale, le moment cinétique $\vec{L}_c$ par rapport au centre C était un vecteur constant. La constante des aires est égale $|\vec{L}_c|/m$.

Pour calculer le moment cinétique, on décompose le vecteur $\vec{CA}$ (figure S31.1) par la relation de Chasles soit $\vec{CA} = \vec{CD} + \vec{DA}$.

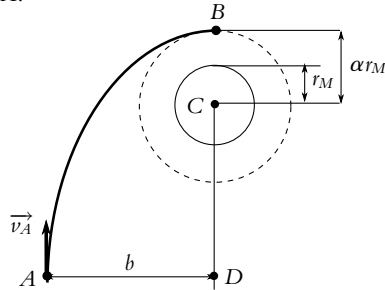


Figure S31.1

Le moment cinétique est donc :

$$\vec{L}_c = \vec{CA} \wedge m \vec{v}_A = m(\vec{CD} \wedge \vec{v}_A + \vec{DA} \wedge \vec{v}_A) = m\vec{DA} \wedge \vec{v}_A = mv_A b \vec{u}_z$$

où $\vec{u}_z$ est le vecteur unitaire perpendiculaire à la figure «entrant» dans la feuille. On en déduit $C = v_A b$.

En B le moment cinétique est $\vec{L}_c = \vec{CB} \wedge m \vec{v}_B$ soit $\vec{L}_c = m\alpha r_M v_B \vec{u}_z$. Par la constance de C on en déduit :

$$C = bv_A = \alpha r_M v_B \Rightarrow v_B = \frac{b}{\alpha r_M} v_A$$

3. On utilise la conservation de l'énergie mécanique :

$$E_m(A) = E_m(B) \Rightarrow \frac{1}{2}mv_A^2 = \frac{1}{2}mv_B^2 - \frac{GM_M m}{r_B} \Rightarrow b = \alpha r_M \sqrt{1 + \frac{2GM_M}{\alpha r_M v_A^2}}$$

L'application numérique pour $\alpha = 3$ donne $b = 15\,603$ km.

4. Le véhicule arriverait juste à la surface de Mars pour $\alpha = 1$ soit :

$$b = r_M \sqrt{1 + \frac{2GM_M}{r_M v_A^2}} = 7\,620 \text{ km}$$

5. Sur une orbite circulaire, on applique la relation fondamentale de la dynamique :

$$-mr_c \dot{\theta}^2 = -\frac{GmM_M}{r_c^2} \Rightarrow v_c = r_c \dot{\theta} = \sqrt{\frac{GM_M}{r_c}}$$

L'application numérique pour $r_c = 3r_M$ donne $v_c = 2,05 \text{ km.s}^{-1}$. La période de révolution est $T = 2\pi r_c / v_c = 8,68 \text{ h}$.

6. Pour passer sur l'orbite circulaire, il faut que la vitesse passe de v_B à v_c soit $\Delta v_c = v_c - v_B$ ou $\Delta v = -1,33 \text{ km.s}^{-1}$.

B.5 1. a) La satellite sur l'orbite géostationnaire a la même vitesse angulaire que la Terre, donc $\Omega_g = 2\pi/T$ avec $T = 24 \text{ h}$ soit $\Omega = 7,3 \cdot 10^{-5} \text{ rad.s}^{-1}$.
Le référentiel d'étude est le référentiel géocentrique $\mathcal{R}_G$.

- b) L'orbite est circulaire ; on applique la relation fondamentale de la dynamique dans $\mathcal{R}_G$:

$$-mR_g\Omega_g^2 = -\frac{GM_T m}{R_g^2} \Rightarrow R_g = \left(\frac{GM_T}{\Omega_g^2}\right)^{1/3}$$

L'application numérique donne $R_g = 42\,180$ km soit une altitude $h_g = R_g - R_T$ et $h = 35\,800$ km.

La vitesse est $v_g = R_g\Omega_g$ soit $v_g = \sqrt{\frac{GM_T}{R_g}}$ ou $v_g = (GM_T\Omega_g)^{1/3}$. L'application numérique donne $v_g = 3,1$ km.s⁻¹.

Le plan de l'orbite contient le centre de la Terre donc pour avoir comme axe l'axe des pôles, l'orbite doit être dans le plan de l'équateur.

2. a) Soit r_p et r_a les rayons du périégée et de l'apogée. On peut écrire $a = (r_p + r_a)/2$ ou $a = (2R_T + h_0 + h_g)/2$ soit $a = 24\,380$ km. D'autre part, on a les égalités $r_p = p/(1+e)$ et $p = a(1-e^2)$. On en déduit $e = 1 - \frac{r_p}{a}$ soit $e = 0,73$ et $p = 11\,390$ km.

- b) La relation fondamentale de la dynamique écrite avec la formule de Binet est :

$$m\vec{a} = -mC^2 \left(\frac{du}{d\theta} + u \right) \vec{u}_r = -GM_T m u^2 \vec{u}_r \Rightarrow \frac{du}{d\theta} + u = \frac{GM_T}{C^2}$$

La solution est :

$$r = \frac{1}{u} = \frac{\frac{C^2}{GM_T}}{1 + A \frac{C^2}{GM_T} \cos \theta}$$

où A est une constante d'intégration. On trouve donc $p = \frac{C^2}{GM_T}$. Le périégée et l'apogée sont les deux seuls points de la trajectoire pour lesquels la vitesse est perpendiculaire au rayon, d'où $C = r_p v_p$ et $C = r_a v_a$. Avec l'expression de C on trouve donc : $v_p = \sqrt{pGM_T}/r_p = 10,3$ km.s⁻¹ et $v_a = \sqrt{pGM_T}/r_a = 1,6$ km.s⁻¹

- c) La troisième loi de Képler $T^2 = a^3 4\pi^2 / (GM_T)$ nous permet de calculer la période sur l'orbite de transfert, soit $T = 38\,800$ s. Le satellite décrit une demi ellipse donc il reste $T/2 = 5,4$ h sur l'orbite de transfert.

B.6 1. a) Le mouvement du point O est circulaire uniforme. La vitesse est obtenue par le produit vectoriel $\vec{v}_O = \vec{\Omega} \wedge \vec{CO}$ où C est le centre de la Terre. On obtient $\vec{v}_O = \Omega R_T \cos \lambda \vec{u}_\varphi$ où $\vec{u}_\varphi$ est le vecteur unitaire tangent au parallèle du lieu, orienté vers l'est.

- b) Pour profiter au maximum de la vitesse donnée par la rotation de la Terre, il faut avoir $\cos \lambda$ le plus grand possible, donc la latitude la plus faible. Parmi les trois propositions, c'est Kourou qui est donc le meilleur site.

- c) Pour calculer l'énergie potentielle, on calcule le travail de la force de gravitation :

$$\delta W = -\frac{mM_T G}{r^2} \vec{u}_r \cdot d\vec{r} = -dE_p \Rightarrow E_p = -\frac{GmM_T}{r} + cte$$

La constante est prise nulle pour avoir une énergie potentielle nulle à l'infini. L'énergie mécanique sur la base de lancement est donc :

$$E_m = E_c + E_p = \frac{1}{2} m \Omega^2 R_T^2 \cos^2 \lambda - \frac{GmM_T}{R_T}$$

- d) On sait que dans le référentiel géocentrique, il faut une énergie mécanique nulle au minimum pour que l'objet se libère de l'attraction terrestre. Il faut donc calculer la vitesse $\vec{v}_G$ dans $\mathcal{R}_G$. La loi de composition des vitesses donne $\vec{v}_G = \vec{v}_l + \vec{v}_e$. La vitesse d'entraînement est celle calculée précédemment et elle est horizontale donc perpendiculaire à $\vec{v}_l$. Cela entraîne $v_G^2 = v_l^2 + v_e^2$, d'où l'énergie mécanique dans $\mathcal{R}_G$:

$$E_m = \frac{1}{2} m (\Omega^2 R_T^2 \cos^2 \lambda + v_l^2) - \frac{GmM_T}{R_T}$$

Pour avoir $E_m = 0$, il faut donc :

$$v_l = \sqrt{\frac{2GM_T}{R_T} - \Omega^2 R_T^2 \cos^2 \lambda} = 11,19 \text{ km.s}^{-1}$$

2. a) On applique le relation fondamentale de la dynamique à une trajectoire circulaire :

$$-mr\dot{\theta}^2 = -\frac{GM_T m}{r^2} \Rightarrow \omega = \dot{\theta} = \sqrt{\frac{GM_T}{r^3}} = \text{cte}$$

Puisque r est constante, la vitesse angulaire est constante et le mouvement est circulaire uniforme.

La vitesse v est égale à $r\dot{\theta}$ soit avec ce qui précède $v = \sqrt{\frac{GM_T}{r}}$.

La période T est égale à $2\pi/\omega$ d'où la troisième loi de Képler :

$$T^2 = 4\pi^2 \frac{r^3}{GM_T} \Rightarrow \frac{T^2}{r^3} = \frac{4\pi^2}{GM_T}$$

- b) Pour l'orbite géostationnaire, la période est celle de rotation de la Terre soit $T = 24$ h et d'après l'expression précédente $r = 42\,290$ km. Le plan de l'orbite contient le centre de la Terre donc pour avoir comme axe l'axe des pôles, l'orbite doit être dans le plan de l'équateur.

- c) le satellite doit avoir au minimum, une énergie mécanique nulle pour échapper à l'attraction terrestre. Initialement, sur son orbite, son énergie mécanique est :

$$E_m = \frac{1}{2}mv^2 - \frac{GM_T m}{R_T + z} = \frac{1}{2}m \frac{GM_T}{R_T + z} - \frac{GM_T m}{R_T + z} = -\frac{1}{2} \frac{GM_T m}{R_T + z}$$

L'énergie cinétique minimale $E_{c,m}$ à lui communiquer est telle que :

$$E_{c,m} + E_m = 0 \Rightarrow E_{c,m} = \frac{1}{2} \frac{GM_T m}{R_T + z}$$

L'application numérique donne $E_{c,m} = 2,8 \cdot 10^{10}$ J cela correspond à une vitesse supplémentaire de 3 km.s^{-1} .

- d) Le satellite est dans un état lié donc l'énergie mécanique est négative, $E_{m0} < 0$. Pour une trajectoire circulaire, on a calculé précédemment l'énergie mécanique et on a vu qu'elle était égale à la moitié de l'énergie potentielle. L'expression montre aussi que $E_m = -E_c$. On suppose que sur des faibles intervalles de temps, la trajectoire reste circulaire donc $dE_m = -dE_c$. Cela paraît paradoxal, mais le freinage augmente l'énergie cinétique donc la vitesse. Comme on va le voir, cela est dû au fait que l'altitude décroît. L'équation précédente donne donc $dE_c = -E_{m0}b \, dt$ soit :

$$E_c - E_{c0} = -E_{m0}bt \Rightarrow v(t) = \sqrt{v_0^2 - \frac{2}{m}E_{m0}bt}$$

Pour le rayon, on utilise l'énergie potentielle avec pour la trajectoire circulaire $E_m = E_p/2$, on a $dE_m = \frac{dE_p}{2}$ et donc $dE_p = \frac{E_{m0}b \, dt}{2}$. En intégrant, on obtient :

$$E_p(t) - E_p(0) = \frac{E_{m0}bt}{2} \Rightarrow r = \left(\frac{1}{r_0} - \frac{E_{m0}bt}{2GM_T m} \right)^{-1}$$

Donc r diminue.

Ainsi $\Delta E_c = -2\Delta E_p$ donc une part de l'énergie mécanique semble perdue. Elle est responsable de l'échauffement du satellite (on l'appellera énergie interne en thermodynamique).

- B.7** 1. a) Le satellite est lié à la Terre, donc sa trajectoire est soit un cercle, soit une ellipse : ici comme r varie, c'est une ellipse.

- b) Pour l'apogée $r_a = p/(1 - e)$ et pour le périégée $r_p = p/(1 + e)$ avec $r_a = R + H$ et $r_p = R + h$. On en tire :

$$e = \frac{H - h}{2R + H + h} = 0,72 \text{ et } p = 2 \frac{(R + h)(R + H)}{2R + H + h} = 11,9 \cdot 10^3 \text{ km}$$

- c) On peut écrire $2a = r_a + r_p = 2R + H + h$, soit $a = 24\,650$ km.

- d) La troisième loi de Képler $\frac{T^2}{a^3} = \frac{4\pi^2}{GM}$ est démontrée dans le cours.

- e) Un satellite géostationnaire a pour période de révolution, la période de rotation de la Terre. Avec la troisième loi de Képler, on obtient :

$$T^2 = \frac{4\pi^2}{GM}(R+H)^3$$

- f) On applique cette loi à l'orbite elliptique d'Hipparcos :

$$T_h^2 = \frac{4\pi^2}{GM}a^3 = T^2 \frac{a^3}{(R+H)^3} = T^2 \left(\frac{2R+h+H}{2(R+H)} \right)^3 \Rightarrow T_h = T \left(\frac{2R+h+H}{2(R+H)} \right)^{3/2}$$

L'application numérique donne $T_h = 10,6$ h.

2. a) On utilise l'équation de l'ellipse en coordonnées polaires $r = \frac{p}{1 - e \cos \theta}$ (le signe «-» est dû au fait qu'on a pris l'origine des angles à l'apogée), ce qui entraîne $\cos \theta = \frac{1}{e} \left(1 - \frac{p}{r} \right)$. On calcule θ pour $r = r_1$ et $r = r_2$ soit $\theta_1 = 125^\circ$ et $\theta_2 = 37^\circ$
- b) On a représenté sur la figure (S31.2) les deux cercles limitant la ceinture de Van Hallen, la trajectoire elliptique du satellite et la zone hachurée correspondant à la surface S_b balayée par la rayon vecteur OM lorsque le satellite passe dans la ceinture de Van Hallen.

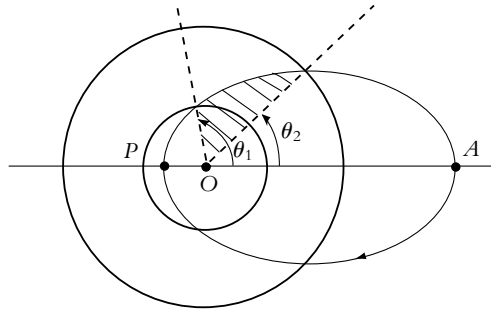


Figure S31.2

- c) D'après la loi des aires, la surface balayée par unité de temps est constante et égale à $C/2$. Hipparcos balaie $2S_b$ par période dans la ceinture de Van Hallen donc :

$$\frac{C}{2} = \frac{S_e}{T_h} = \frac{2S_b}{t_0} \Rightarrow \rho = \frac{t_0}{T_h} = \frac{2S_b}{\pi ab} = 0.44$$

- B.8** 1. On étudie le satellite dans le référentiel géocentrique supposé galiléen. Il est soumis à la force

$\vec{f} = -\frac{GM_T m}{r^2} \vec{u}_r$ qui est une force centrale. Par conséquent, le moment de cette force par rapport au centre de la Terre est nul et le moment cinétique est constant. Le mouvement a donc lieu dans le plan contenant le centre de la Terre.

Comme le satellite est géostationnaire, il suit la Terre dans sa rotation sur elle-même donc son mouvement a lieu dans un plan perpendiculaire à l'axe de rotation de la Terre sur elle-même.

Le seul plan à la fois perpendiculaire à l'axe de rotation de la Terre sur elle-même et passant par son centre est le plan équatorial. C'est donc dans ce plan qu'a lieu le mouvement d'un satellite géostationnaire.

2. Le mouvement est plan pour chaque orbite ainsi que pour la trajectoire de transfert. Les conditions initiales d'une orbite sont les conditions finales de la trajectoire précédente donc les trois orbites évoluent dans un même plan qui est forcément le plan équatorial puisque le mouvement final est géostationnaire.
3. On projette le principe fondamental de la dynamique appliqué au satellite sur $\vec{u}_r$ soit $-mr\ddot{\theta} = -\frac{GM_T m}{r^2}$. Or le module de la vitesse s'écrit pour une trajectoire circulaire ($\dot{r} = 0$) s'écrit

$v = r\dot{\theta}$ soit en reportant dans l'équation précédente $r\frac{v^2}{r^2} = \frac{GM_T}{r^2}$ donc $v = \frac{GM_T}{r}$ et dans le cas de l'orbite initiale à la surface de la Terre $v_B = \sqrt{\frac{GM_T}{R_T}}$.

4. L'application numérique donne $v_B = 7,91.10^3 \text{ m.s}^{-1} = 28,5.10^3 \text{ km.h}^{-1}$.

5. On applique la troisième loi de Kepler soit $\frac{T^2}{R_G^3} = \frac{4\pi^2}{GM_T}$. On en déduit $R_G = \sqrt[3]{\frac{GM_T T^2}{4\pi^2}}$.

6. L'application numérique donne $R_G = 42,2.10^3 \text{ km}$.

7. Pour obtenir l'altitude, il suffit de retrancher au rayon qui vient d'être trouvé le rayon terrestre soit $h = R_G - R_T = 35,8.10^3 \text{ km}$.

8. On a la même relation qu'à la question 3 avec $r = R_G$ soit

$$v_G = \sqrt{\frac{GM_T}{R_G}} = 3,07.10^3 \text{ m.s}^{-1} = 11,1.10^3 \text{ km.h}^{-1}$$

9. Les points A et P correspondent respectivement à l'apogée et au périégée. En ces points, la vitesse est orthoradiale car on a un extremum de la distance au centre de la Terre autrement dit r est maximal ou $\dot{r} = 0$, cette dernière relation impliquant que la composante radiale de la vitesse est nulle. On en déduit qu'en ces points $v = v_\theta = r\dot{\theta}$ et en utilisant la constante des aires, on en déduit $C = r^2\dot{\theta} = r_A v_A = r_P v_P$.

10. L'énergie mécanique Em est conservée. En exprimant l'énergie mécanique à l'apogée et au périégée, on pourra écrire l'égalité des deux expressions obtenues. A l'apogée, on a $Em = \frac{1}{2}mv_A^2 - \frac{GM_T m}{r_A}$ et au périégée $Em = \frac{1}{2}mv_P^2 - \frac{GM_T m}{r_P}$. En multipliant la première relation par r_A^2 et la seconde par r_P^2 avant de les soustraire, on obtient $(r_A^2 - r_P^2) Em = \frac{1}{2}m(r_A^2 v_A^2 - r_P^2 v_P^2) - GM_T m(r_A - r_P)$ avec la conservation de l'énergie mécanique. La relation établie à la question précédente permet d'obtenir $r_A^2 v_A^2 - r_P^2 v_P^2 = 0$ et en factorisant $r_A^2 - r_P^2 = (r_A - r_P)(r_A + r_P)$, on en déduit $(r_A - r_P)(r_A + r_P) Em = -GM_T m(r_A - r_P)$. Or $r_A + r_P = 2a$ donc en simplifiant par $r_A - r_P \neq 0$, on a l'expression demandée $Em = -\frac{GM_T m}{2a}$.

11. On utilise le fait que $2a = R_G + R_T$ pour déterminer la valeur numérique de l'énergie mécanique $Em = -1,23.10^{10} \text{ J}$.

12. L'énergie mécanique s'écrit dans le cas général

$$Em = \frac{1}{2}mv^2 - \frac{GM_T m}{r} = -\frac{GM_T m}{R_G + R_T}$$

par conservation de l'énergie mécanique. On en déduit

$$v = \sqrt{2GM_T \left(\frac{1}{r} - \frac{1}{R_G + R_T} \right)}$$

13. Au périégée P , il faut imposer une variation de vitesse

$$\Delta v_P = v_P - v_B = \sqrt{2GM_T \left(\frac{1}{R_T} - \frac{1}{R_T + R_G} \right)} - \sqrt{\frac{GM_T}{R_T}}$$

soit numériquement $\Delta v_P = 2,50.10^3 \text{ m.s}^{-1} = 8,96.10^3 \text{ km.h}^{-1}$.

14. La variation d'énergie mécanique s'écrit

$$\Delta Em_P = Em - Em_B = -GM_T m \left(\frac{1}{R_G + R_T} - \frac{1}{2R_T} \right) = -3,46.10^{10} \text{ J}$$

15. De même, à l'apogée, on impose une variation de vitesse

$$\Delta v_A = v_G - v_A = \sqrt{\frac{GM_T}{R_T}} - \sqrt{2GM_T \left(\frac{1}{R_G} - \frac{1}{R_T + R_G} \right)}$$

soit numériquement $\Delta v_A = 1,50 \cdot 10^3 \text{ m.s}^{-1} = 5,4 \cdot 10^3 \text{ km.h}^{-1}$.

16. La variation d'énergie mécanique s'écrit

$$\Delta E_{m_A} = E_{m_G} - E_m = -GM_T m \left(\frac{1}{2R_G} - \frac{1}{R_G + R_T} \right) = -5,23 \cdot 10^{10} \text{ J}$$

17. La troisième loi de Kepler s'écrit $\frac{T^2}{a^3} = \frac{4\pi^2}{GM_T}$. Or $a = \frac{R_G + R_T}{2}$, on en déduit la période

$T = \pi \sqrt{\frac{(R_G + R_T)^3}{2GM_T}}$ de la trajectoire de transfert. Le transfert a une durée égale à la moitié de la période soit par l'application numérique 5 h 13 min 45 s.

- B.9** 1. On étudie la Lune dans le référentiel terrestre $\mathcal{R}_0$ supposé galiléen. L'astre est soumis à la force d'attraction terrestre uniquement. La projection du principe fondamental de la dynamique sur Ox_1 s'écrit $-M\omega_1^2 = -\frac{GMkM}{a^2}$. On en déduit $\omega_1^2 = \frac{GkM}{a^3}$.

2. On cherche l'équigravité c'est-à-dire que les deux forces de gravitation se compensent soit $\frac{GkMm}{(x-a)^2} = \frac{GMm}{x^2}$. On en déduit l'équation $x^2(k-1) + 2ax - a^2 = 0$. Son discriminant vaut

$\Delta = 4a^2 + 4a^2(k-1) = 4a^2k > 0$. Les solutions sont donc $x = \frac{-2a \pm 2a\sqrt{k}}{2(k-1)}$, la seule solution acceptable est la solution positive soit $x = \frac{a}{1 + \sqrt{k}}$.

3. L'application numérique donne $x = 38,4 \cdot 10^6 \text{ m}$.

4. L'égalité des vitesses de rotation de la Lune sur elle-même et autour de la Terre montre que la Lune présente toujours la même face à la Terre autrement dit qu'on voit toujours la même face de la Lune.

5. Un satellite est sélénostationnaire si sa vitesse de rotation autour de la Lune est égale à la vitesse de rotation de la Lune sur elle-même soit $\omega_1 = \omega_2$.

6. On étudie le satellite dans $\mathcal{R}_1$ qui n'est pas galiléen puisque c'est le référentiel $\mathcal{R}_0$ qui est supposé galiléen. Le satellite est donc soumis à l'attraction de la Lune, à la force d'inertie d'entraînement $m\omega_1^2 \vec{TS}$ et à la force d'inertie de Coriolis $2m\omega_1^2 \vec{LS}$. On note qu'il est soumis également à l'attraction terrestre mais l'énoncé demande de négliger cette force. Le principe fondamental de la dynamique s'écrit alors $-m\omega_1^2 \vec{LS} = -\frac{GMm}{LS^3} \vec{LS} + m\omega_1^2 \vec{TL} + m\omega_1^2 \vec{LS} + 2m\omega_1^2 \vec{LS}$. On a donc à résoudre après simplification l'équation $\omega_1^2(4r+a) = \frac{GM}{r^2}$. En négligeant $4r$ devant a , on en déduit $r = \frac{a}{\sqrt{k}}$.

7. L'application numérique donne $r = 41,7 \cdot 10^6 \text{ m}$.

8. La position est proche de celle d'équigravité donc on ne peut pas négliger l'attraction terrestre.

9. On procède comme à la question 5 mais en tenant compte de la gravitation terrestre c'est-à-dire en ajoutant au bilan des forces la force d'attraction terrestre soit $-G\frac{kMm}{TS^3} \vec{TS}$.

On en déduit les deux équations

$$\omega_1^2 xq = \frac{GM}{x^2 a^2} + \omega_1^2 a - \omega_1^2 ax - 2\omega_1^2 ax - \frac{GkM}{a^2(1-x)^2}$$

et

$$-\omega_1^2 xq = -\frac{GM}{x^2 a^2} + \omega_1^2 a + \omega_1^2 ax + 2\omega_1^2 ax - \frac{GkM}{a^2(1-x)^2}$$

soit
$$GM \left(\frac{k}{a^2 (1 \pm x)^2} \pm \frac{1}{a^2} \right) = \omega_1^2 a (1 \pm 4x)$$

10. On effectue un développement limité avec $x \ll 1$. Après calcul, on obtient $x = \pm \frac{1}{6k}$.

11. L'application numérique donne une distance $LP = \pm \frac{a}{6k} = \pm 7,91 \cdot 10^5$ m.

12. En appliquant la troisième loi de Kepler, on a $\frac{T_0^2}{a^3} = \frac{T^2}{R_T^3}$ soit

$$\tau = \frac{T}{2} = \frac{T_0}{2} \sqrt{\left(\frac{R_T}{a} \right)^3}$$

13. L'application numérique donne $\tau = 42$ min 16 s.

Chapitre 32

A.1 Le système, constitué de la valise, étudié dans le référentiel associé à l'ascenseur, est soumis à son poids, à la force d'inertie d'entraînement due à l'accélération de l'ascenseur par rapport au référentiel terrestre et à la force exercée par le voyageur. On notera que la force d'inertie d'entraînement n'est pas incluse dans le poids : celle qui fait partie du poids tient compte du caractère non galiléen du référentiel géocentrique (ou ici terrestre, ce qui revient à ne pas considérer l'influence de la rotation de la Terre sur elle-même) par rapport au référentiel de Copernic. Si on se place dans un référentiel qui subit une accélération par rapport au référentiel géocentrique (ou terrestre), il est nécessaire de considérer la nouvelle force d'inertie d'entraînement associée.

La valise étant immobile par rapport à l'ascenseur, on détermine la force exercée par le voyageur sur la valise : $\vec{T} = m(\vec{a}_e - \vec{g})$ en appliquant le principe fondamental de la dynamique. Il s'agit de l'opposée de la force qu'exerce la valise sur le voyageur par le principe des actions réciproques.

1. Quand l'ascenseur démarre en montant, $\vec{a}_e$ est dirigée vers le haut donc la norme de $\vec{T}$ augmente et la valise paraît plus lourde.
2. Quand l'ascenseur démarre en descendant, $\vec{a}_e$ est dirigée vers le bas donc la norme de $\vec{T}$ diminue et la valise paraît moins lourde.
3. Quand l'ascenseur est arrêté, $\vec{a}_e$ est nulle donc la norme de $\vec{T}$ est égale à celle du poids de la valise qui semble conserver son poids.
4. Quand l'ascenseur ralentit en montant, $\vec{a}_e$ est dirigée vers le bas donc la norme de $\vec{T}$ diminue et la valise paraît moins lourde.
5. Quand l'ascenseur ralentit en descendant, $\vec{a}_e$ est dirigée vers le haut donc la norme de $\vec{T}$ augmente et la valise paraît plus lourde.

A.2 1. Le point matériel M est soumis dans le référentiel $\mathcal{R}$:

- à l'attraction terrestre $\vec{A}_T(M) = -\frac{GM_T m}{OM^3} \vec{OM}$,
- à la force d'inertie d'entraînement $\vec{f}_{ie} = m\omega_0^2 \vec{HM}$

En effet, le mouvement du point coïncidant est un mouvement circulaire de centre H , projection de M sur l'axe OZ .

- à la force d'inertie de Coriolis $\vec{f}_{ic} = -2m\omega_0 \vec{u}_z \wedge \vec{v}_{/\mathcal{R}}(M)$.

Le principe fondamental s'écrit alors :

$$m\vec{a}_{/\mathcal{R}}(M) = \frac{GM_T m}{OM^3} \vec{OM} + m\omega_0^2 \vec{HM} - 2m\omega_0 \vec{u}_z \wedge \vec{v}_{/\mathcal{R}}(M) \quad (1)$$

On explicite ces forces dans la base $(\vec{u}_x, \vec{u}_y, \vec{u}_z)$ et on effectue un développement limité de l'attraction terrestre au premier ordre en $\frac{x}{R}$, $\frac{y}{R}$ et $\frac{z}{R}$. On obtient :

$$\begin{aligned}\vec{A}_T(M) &= -\frac{GM_T m}{R^3} \left(1 - 3\frac{x}{R}\right) ((x+R)\vec{u}_x + y\vec{u}_y + z\vec{u}_z) \\ &= -\frac{GM_T m}{R^2} \left(1 - 3\frac{x}{R}\right) \vec{u}_x - \frac{GM_T m}{R^3} (x\vec{u}_x + y\vec{u}_y + z\vec{u}_z)\end{aligned}$$

puis

$$\vec{f}_{ic} = m\omega_0^2 ((R+x)\vec{u}_x + y\vec{u}_y) \quad \text{et} \quad \vec{f}_{ic} = 2m\omega_0^2 (y\vec{u}_x - x\vec{u}_y)$$

Sachant que le satellite a une trajectoire circulaire dans le référentiel $\mathcal{R}_0$, le principe fondamental de la dynamique appliqué au satellite dans ce référentiel s'écrit :

$$m_{\text{sat}} \vec{a}/\mathcal{R}_0(S) = -\frac{GM_T m_{\text{sat}}}{R^3} \vec{OS} = -m_{\text{sat}} R\omega_0^2 \vec{u}_x. \text{ On en déduit : } \omega_0^2 = \frac{GM_T}{R^3}.$$

En reportant toutes ces expressions dans l'équation (1), on trouve bien le système proposé.

2. La résolution du système (en utilisant la méthode de substitution) donne, compte tenu des conditions initiales :

$$\begin{cases} x(t) = 4x_0 - 3x_0 \cos \omega_0 t \\ y(t) = y_0 + 6x_0 \sin \omega_0 t - 6\omega_0 x_0 t \\ z(t) = z_0 \cos \omega_0 t \end{cases} \quad \text{puis} \quad \begin{cases} \dot{x}(t) = 3x_0 \omega_0 \sin \omega_0 t \\ \dot{y}(t) = 6x_0 \omega_0 \cos \omega_0 t - 6\omega_0 x_0 \\ \dot{z}(t) = -z_0 \omega_0 \sin \omega_0 t \end{cases}$$

On note qu'ici la méthode utilisant la variable complexe $u = x + iy$ n'est pas possible : les dérivées premières et secondes par rapport au temps se combinent mais il reste un terme en x qui n'a pas d'équivalent en y .

Le mouvement est donc la combinaison d'un mouvement d'oscillation à la pulsation ω_0 sur chaque axe et d'une translation selon Sy . On a donc une dérive suivant cet axe. La vitesse de dérive correspondante est $\vec{v}_D = \langle \vec{v} \rangle = -6\omega_0 x_0 \vec{u}_y$.

3. Pour accéder à ω_0 , on peut mesurer soit la vitesse de dérive proportionnelle à ω_0 soit directement la période du mouvement d'oscillation selon Oz ou dans le plan xOy .

A.3 1. La période de rotation de la Terre dans le référentiel géocentrique est en moyenne $T = 24$ h donc sa vitesse angulaire est $\Omega = 2\pi/T$ soit $\Omega = 7,3 \cdot 10^{-5} \text{ rad.s}^{-1}$.

2. La projection de $\vec{\Omega}$ se déduit de la figure (S32.1), ce qui donne :

$$\vec{\Omega} = \Omega \cos L \vec{u}_y + \Omega \sin L \vec{u}_z$$

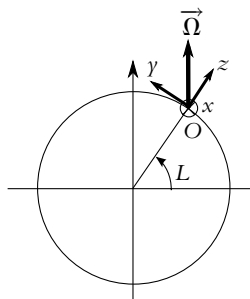


Figure S32.1

3. Le référentiel $\mathcal{R}_T$ est en rotation par rapport au référentiel géocentrique galiléen, donc il faut tenir compte des forces d'inertie d'entraînement et de coriolis. Il faut cependant se rappeler que la force d'inertie d'entraînement est incluse dans le poids par définition du poids. Les forces s'exerçant sur le wagon sont :

- le poids $m\vec{g}$,

- la réaction normale des rails $\vec{F} = F_H \vec{u}_Y + F_V \vec{u}_Z$,
 - la force d'inertie de coriolis $\vec{F}_{ic}$,
 - la force motrice $\vec{F}_m$ équilibrée par la force de frottement $\vec{F}_f$ colinéaires aux rails.
- La relation fondamentale de la dynamique dans $\mathcal{R}_T$:

$$m \vec{a} / \mathcal{R}_T = m \vec{g} + \vec{F} + \vec{F}_{ic} + \vec{F}_m + \vec{F}_f$$

On a représenté les différents repères sur la figure (S32.2).

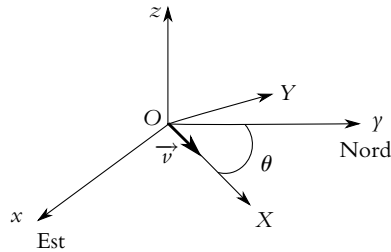


Figure S32.2

Dans le repère (O, X, Y, Z) le vecteur rotation s'écrit :

$$\vec{\Omega} = \Omega(\cos L \cos \theta \vec{u}_X + \cos L \sin \theta \vec{u}_Y + \sin L \vec{u}_Z)$$

Avec la vitesse égale à $\vec{v} / \mathcal{R}_T = v \vec{u}_X$, la force d'inertie de coriolis est :

$$\vec{F}_{ic} = -2M\vec{\Omega} \wedge \vec{v} / \mathcal{R}_T = -2m\Omega v(\sin L \vec{u}_Y - \cos L \sin \theta \vec{u}_Z)$$

On projette la relation fondamentale sur le repère (O, X, Y, Z) sachant que l'accélération est nulle (vitesse constante) :

$$\begin{cases} 0 = F_m + F_f \\ 0 = F_H - 2M\Omega v \sin L \\ 0 = -Mg + F_V + 2M\Omega v \cos L \sin \theta \end{cases}$$

4. On en déduit :

$$F_H = 2Mv\Omega \sin L \text{ et } F_V = mg - 2Mv\Omega \cos L \sin \theta$$

soit pour $\theta = \pi/2$:

$$F_H = 2Mv\Omega \sin L \text{ et } F_V = Mg - 2Mv\Omega \cos L$$

L'application numérique donne $F_H = 82 \text{ N}$.

5. La force horizontale exercée par le wagon sur les rails est $-F_H$, ce qui revient à exercer une force plus importante sur le rail gauche par rapport au sens de déplacement du train. Ainsi, pour les voies à sens unique de circulation le rail gauche s'use plus vite que le rail droit.

B.1 1. On étudie la fusée dans le référentiel terrestre non galiléen. Elle est soumise au poids incluant la force d'inertie d'entraînement et à la force d'inertie de Coriolis. La traduction du principe fondamentale la dynamique dans la base $Oxyz$ donne $\ddot{x} = 2\Omega(\dot{y} \sin \lambda - \dot{z} \cos \lambda)$, $\ddot{y} = -2\Omega\dot{x} \sin \lambda$ et $\ddot{z} = -g + 2\Omega\dot{x} \cos \lambda$.

2. Si on ne tient pas compte de la force d'inertie de Coriolis, on intègre deux fois par rapport au temps $\ddot{x} = 0$, $\ddot{y} = 0$ et $\ddot{z} = -g$ avec les conditions initiales $\vec{v}(0) = v_0 \vec{u}_z$ et $x(0) = y(0) = z(0) = 0$. On obtient $x = 0$, $y = 0$ et $z = \frac{-gt^2}{2} + v_0 t$.

3. L'altitude maximale est obtenue pour $\dot{z} = 0$ c'est-à-dire à l'instant $t = \frac{v_0}{g}$. L'altitude obtenue est

$$\text{donc } h = \frac{v_0^2}{2g}.$$

4. En effectuant l'approximation à l'ordre 0 proposée, on a $\dot{x} = 0$, $\dot{y} = 0$ et $\dot{z} = -gt + v_0$. En reportant dans les équations du mouvement, on obtient $\ddot{x} = -2\Omega \cos \lambda (v_0 - gt)$, $\ddot{y} = 0$ et $\ddot{z} = -g$. On intègre à nouveau deux fois de suite avec les mêmes conditions initiales qu'en 2. On en déduit $x = \Omega \cos \lambda \left(\frac{gt^3}{3} - v_0 t^2 \right)$, $y = 0$ et $z = -\frac{1}{2}gt^2 + v_0 t$.
5. Le retour au sol correspond à $z = 0$. La résolution de $z(t) = 0$ donne $t = 0$ qui correspond à l'instant initial et qui doit donc être exclu ou $t = \frac{2v_0}{g}$. Dans ce dernier cas, on a $x = -\frac{4\Omega v_0^3 \cos \lambda}{3g^2}$.
6. La valeur de x est toujours négative donc la déviation a lieu vers l'ouest.
7. L'application numérique donne $\Omega = 7,27 \cdot 10^{-5} \text{ rad.s}^{-1}$ et $v_0 = 13,9 \text{ m.s}^{-1}$. On en déduit la hauteur $h = 9,8 \text{ m}$ et la déviation $x = 1,9 \text{ mm}$.
8. D'après la question 5, on en déduit que pour obtenir une déviation donnée, il faut lancer la fusée à la vitesse $v_0 = \sqrt[3]{\frac{3g^2 x}{4\Omega \cos \lambda}} = 507 \text{ km.h}^{-1}$ et alors l'altitude maximale est $h = 1,01 \text{ km}$.

B.2 1. On se réfère à la figure (S32.3) pour les projections ce qui donne :

$$\vec{\Omega} = \Omega \cos \lambda \vec{u}_y + \Omega \sin \lambda \vec{u}_z$$

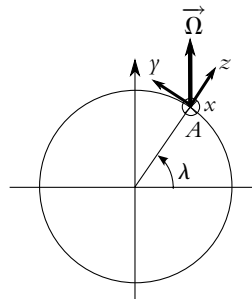


Figure S32.3

2. On écrit la relation fondamentale de la dynamique pour le point M dans le référentiel terrestre. On rappelle que la force d'inertie d'entraînement est incluse dans le poids. On obtient donc :

$$m \vec{a} / \mathcal{R}_T = m \vec{g} + \vec{F}_{ic}$$

Sachant que $\vec{v} / \mathcal{R}_T = \dot{x} \vec{u}_x + \dot{y} \vec{u}_y + \dot{z} \vec{u}_z$, la force d'inertie de coriolis est égale à :

$$\vec{F}_{ic} = -2m \vec{\Omega} \wedge \vec{v} / \mathcal{R}_T = -2m \Omega ((\dot{z} \cos \lambda - \dot{y} \sin \lambda) \vec{u}_x + \dot{x} \sin \lambda \vec{u}_y - \dot{x} \cos \lambda \vec{u}_z)$$

La relation fondamentale projetée sur les trois axes s'écrit :

$$\begin{cases} m\ddot{x} = -2m\Omega \dot{z} \cos \lambda + 2m\Omega \dot{y} \sin \lambda \\ m\ddot{y} = -2m\Omega \dot{x} \sin \lambda \\ m\ddot{z} = -mg + 2m\Omega \dot{x} \cos \lambda \end{cases}$$

3. On intègre l'équation en $\ddot{z}$:

$$\dot{z} = -gt + 2\Omega x \cos \lambda + cte$$

Sachant qu'à $t = 0$ $x = 0$ et $\dot{z} = v_0 \sin \alpha$, la constante est égale à $v_0 \sin \alpha$. On intègre l'équation en $\ddot{y}$:

$$\dot{y} = -2\Omega x \sin \lambda + cte$$

Sachant qu'à $t = 0$ $x = 0$ et $\dot{y} = v_0 \cos \alpha$, la constante est égale à $v_0 \cos \alpha$. On reporte dans l'équation en $\ddot{x}$ et on trouve l'équation demandée.

4. Estimation des différents termes (en m.s^{-2}) :

- $4\Omega^2 x \simeq 4 \times 49.10^{-10} \times 10^2 \simeq 2.10^{-6}$,
- $\Omega g t \simeq 7.10^{-5} \times 10 \times 60 \simeq 4.10^{-2}$,
- $\Omega v_0 \simeq 7.10^{-5} \times 10^3 \simeq 7.10^{-2}$.

On garde donc les deux derniers termes d'où l'équation :

$$\ddot{x} = 2\Omega g t \cos \lambda - 2\Omega v_0 \cos \lambda \sin \alpha + 2\Omega v_0 \cos \alpha \sin \lambda = 2\Omega g t \cos \lambda + 2\Omega v_0 \sin(\lambda - \alpha)$$

On intègre deux fois cette équation, sachant que $\dot{x}(0) = 0$ et $x(0) = 0$, ce qui donne :

$$x = \frac{1}{3}\Omega g t^3 \cos \lambda + \Omega v_0 t^2 \sin(\lambda - \alpha)$$

5. On estime l'ordre de grandeur de chaque terme dans l'expression de $\dot{z}$:

- $g t \simeq 10 \times 60 \simeq 600 \text{m.s}^{-1}$,
- $\Omega x \simeq 7.10^{-5} \times 100 \simeq 7.10^{-3} \text{m.s}^{-1}$,
- $v_0 \sin \alpha \simeq 1000 \text{m.s}^{-1}$.

On néglige donc le second terme d'où :

$$\dot{z} = -g t + v_0 \sin \alpha \Rightarrow z = -\frac{g t^2}{2} + v_0 t \sin \alpha$$

car $z(0) = 0$.

L'instant de retour au sol est obtenu pour $z(t_s) = 0$ soit $t_s = \frac{2v_0 \sin \alpha}{g}$. L'application numérique donne $t_s = 144 \text{ s}$.

6. On reporte l'expression de t_s dans l'expression de x ce qui donne la déviation selon Ax lorsque le projectile touche le sol :

$$x_s = \frac{4\Omega v_0^3 \sin^2 \alpha}{g^2} \left(\frac{2}{3} \sin \alpha \cos \lambda + \sin(\lambda - \alpha) \right) = 323 \text{ m}$$

La déviation a bien lieu vers la droite tant que $\sin(\lambda - \alpha) + \frac{2}{3} \sin \alpha \cos \lambda > 0$ soit tant que $\alpha < 65^\circ$ ce qui est le cas en général pour un tir.

Pour un tir horizontal ($\alpha = 0$) et :

$$\ddot{x} = 2\Omega g t \cos \lambda + 2\Omega v_0 \sin \lambda \Rightarrow x = \frac{1}{3}\Omega g t^3 \cos \lambda + \Omega v_0 t^2 \sin \lambda$$

Il est difficile de conclure, mais c'est en général le terme en $\Omega v_0 t^2$ qui l'emporte sur un temps raisonnable donc on a une déviation vers la droite dans l'hémisphère nord ($\sin \lambda > 0$) et vers la gauche dans l'hémisphère sud ($\sin \lambda < 0$).

- B.3** 1. a) Dans le référentiel $\mathcal{K}$, on applique la relation fondamentale de la dynamique à la cabine spatiale soumise uniquement à l'attraction terrestre :

$$-m(R_T + h)\omega^2 = -\frac{GM_T m}{(R_T + h)^2}$$

On connaît l'accélération de la pesanteur à la surface de la Terre soit $g_0 = \frac{GM_T}{R_T^2}$. On en déduit :

$$\omega = \sqrt{\frac{g_0 R_T^2}{(R_T + h)^3}}$$

b) La période est $T = 2\pi/\omega$ soit $T = 5560 \text{ s}$ ou $1,54 \text{ h}$.

2. Le référentiel $\mathcal{K}'$ n'est pas galiléen car son mouvement par rapport à $\mathcal{K}$ n'est pas rectiligne uniforme.
3. Le point coïncident au point M à l'instant t est le point P fixe dans $\mathcal{K}'$ qui coïncide avec M à l'instant t . Il a un mouvement circulaire dans $\mathcal{K}$ et son accélération est $\vec{a}_e = -\omega^2 \vec{OM}$ ou $\vec{a}_e = -\omega^2(\vec{R} + \vec{r})$. La force d'inertie d'entraînement est donc $\vec{f}_{ie} = -m\vec{a}_e = m\omega^2(\vec{R} + \vec{r})$.
4. Si la particule est mobile dans $\mathcal{K}'$, il faut tenir compte de la force d'inertie de coriolis $\vec{f}_{ic} = -2m\vec{\omega} \wedge \vec{v}$.

5. La force d'attraction gravitationnelle sur M s'écrit :

$$\vec{F} = -\frac{GmM_T}{\|\vec{R} + \vec{r}\|^3}(\vec{R} + \vec{r})$$

On effectue un développement limité du dénominateur à l'ordre 1 :

$$\|\vec{R} + \vec{r}\|^{-3} = (R^2 + 2\vec{R} \cdot \vec{r} + r^2)^{-3/2} = \frac{1}{R^3} \left(1 + \frac{2R_x}{R^2} + \frac{r^2}{R^2}\right)^{-3/2} \simeq \frac{1}{R^3} \left(1 - \frac{3x}{R}\right)$$

Avec $\omega = \sqrt{GM_T/R^3}$, la force s'écrit :

$$\vec{F} = -m\omega^2(\vec{R} + \vec{r} - \frac{3x}{R}(\vec{R} + \vec{r})) \simeq -m\omega^2(\vec{R} + \vec{r} - \frac{3x}{R}\vec{R}) \simeq -m\omega^2(\vec{R} + \vec{r} - 3xu_x)$$

6. La relation fondamentale de la dynamique dans $\mathcal{K}'$ s'écrit :

$$m \frac{d\vec{a}/\mathcal{K}'}{dt} = \vec{F} + \vec{f}_{ic} + \vec{f}_{\text{cor}}$$

L'énoncé fait l'hypothèse que le mouvement a lieu dans le plan xSy donc la force d'inertie de coriolis est :

$$\vec{f}_{ic} = -2m\omega \vec{u}_z \wedge (\dot{x}\vec{u}_x + \dot{y}\vec{u}_y) = -2m\omega(-\dot{y}\vec{u}_x + \dot{x}\vec{u}_y)$$

On peut donc écrire la projection de la relation fondamentale sur Sx et Sy ce qui donne après simplification :

$$\begin{cases} \ddot{x} = 3\omega^2 x + 2\omega \dot{y} \\ \ddot{y} = -2\omega \dot{x} \end{cases}$$

7. On intègre la deuxième équation, ce qui donne $\dot{y} = -2\omega x + cte$. À $t = 0$, $\dot{y} = 0$ et $x = 0$ donc la constante est nulle. On reporte le résultat dans la première équation :

$$\ddot{x} = 3\omega^2 x - 4\omega^2 x \Rightarrow \ddot{x} + \omega^2 x = 0 \Rightarrow x = A \cos \omega t + B \sin \omega t$$

Les conditions initiales $x(0) = 0$ et $\dot{x}(0) = -v_0$ entraînent $A = 0$ et $B = -v_0/\omega$. On reporte ce résultat dans l'expression de $\dot{y}$, ce qui donne :

$$\dot{y} = 2v_0 \sin \omega t \Rightarrow y = 2\frac{v_0}{\omega}(1 - \cos \omega t)$$

avec la condition initiale $y(0) = 0$. En combinant les deux équations, on trouve l'équation d'une ellipse :

$$\left(\frac{x\omega}{v_0}\right)^2 + \left(\frac{y\omega}{2v_0} - 1\right)^2 = 1$$

La période est $T = 2\pi/\omega$.

- B.4** 1. Le système constitué de la masse m dans le référentiel terrestre considéré comme galiléen est soumis à son poids et à la tension $\vec{T}$ du filin.

a) Les deux forces sont coplanaires et comme il n'y a pas de vitesse initiale, le mouvement restera dans le plan vertical défini par la verticale et la direction initiale du filin.

b) On applique le théorème du moment cinétique en notant θ l'angle entre le filin et la verticale. Le moment cinétique par rapport à A s'écrit : $\vec{L}_A = \vec{AM} \wedge m\vec{v} = ml^2 \dot{\theta} \vec{u}_x$. Le moment de la tension du fil est nul puisque la force est colinéaire à $\vec{AM}$ et celui du poids vaut : $\vec{M}_A(m\vec{g}) = -mgl \sin \theta \vec{u}_x$. L'expression du théorème du moment cinétique conduit à : $\ddot{\theta} + \frac{g}{l} \sin \theta = 0$. Comme α est l'angle maximal d'oscillation et qu'il reste faible, on peut approximer $\sin \theta$ par θ et ainsi linéariser cette équation : $\ddot{\theta} + \frac{g}{l} \theta = 0$. On a donc des oscillations à la pulsation $\omega_0 = \sqrt{\frac{g}{l}}$: $\theta = \alpha \cos \omega_0 t$.

c) L'angle θ reste faible donc en effectuant des développements limités, on obtient l'expression de l'altitude $z = l(1 - \cos \theta) \simeq \frac{l\theta^2}{2}$ et celle du déplacement horizontal $d = l \sin \theta \simeq l\theta$.

L'amplitude des oscillations selon la verticale est $z_{\max} = \frac{l\alpha^2}{2} = 25,5 \text{ cm}$ et celle sur l'horizontale $d_{\max} = l\alpha = 5,85 \text{ m}$. On n'oubliera pas d'exprimer les angles en radians pour l'application numérique.

Les vitesses valent : $\dot{z} = l\theta\dot{\theta} = -\omega_0 l\alpha^2 \cos \omega_0 t \sin \omega_0 t$ et $\dot{d} = l\dot{\theta} = -\omega_0 l\alpha \sin \omega_0 t$. Leurs amplitudes sont $\dot{z}_{\max} = \omega_0 l\alpha^2 = 19,5 \text{ cm.s}^{-1}$ et $\dot{d}_{\max} = \omega_0 l\alpha = 2,24 \text{ m.s}^{-1}$ en prenant $g = 9,81 \text{ m.s}^{-2}$.

Quant aux accélérations, elles valent $\ddot{z} = -l\alpha^2 \omega_0^2 \cos 2\omega_0 t$ et $\ddot{d} = -\omega_0^2 l\alpha \cos \omega_0 t$. Leurs amplitudes sont $\ddot{z}_{\max} = l\alpha^2 \omega_0^2 = 7,47 \text{ cm.s}^{-2}$ et $\ddot{d}_{\max} = \omega_0^2 l\alpha = 85,6 \text{ cm.s}^{-2}$.

- La Terre tourne sur elle-même en $T = 24 \text{ h}$ soit une vitesse angulaire $\Omega = \frac{2\pi}{T} = 7,27 \cdot 10^{-5} \text{ rad.s}^{-1}$.
- Le fait de tenir compte de la rotation de la Terre sur elle-même conduit à considérer la force d'inertie de Coriolis (la force d'entraînement est déjà prise en compte dans le poids). On cherche donc à comparer les normes des différentes forces en présence. La force d'inertie de Coriolis a pour expression : $-2m\vec{\Omega} \wedge \vec{v}$, l'ordre de grandeur de son module est : $2m\Omega v$. Cette force sera négligeable devant le poids tant que $2m\Omega v \ll mg$ soit pour une vitesse $v \ll \frac{g}{2\Omega} = 67,5 \cdot 10^3 \text{ m.s}^{-1} = 243 \cdot 10^3 \text{ km.h}^{-1}$. De telles vitesses ne seront pas atteintes par ce pendule. On peut donc considérer que la force d'inertie de Coriolis constitue une correction au poids. On va donc chercher une correction au mouvement précédemment établi.
- Le fait d'écarter la masse dans le plan méridien revient à lui donner à la latitude de l'équateur une vitesse quasi-colinéaire au vecteur rotation de la Terre sur elle-même. Par conséquent, la force d'inertie de Coriolis est quasi-nulle au niveau de l'équateur pour ce système qui ne pourra donc détecter la rotation de la Terre sur elle-même.
- On remarque que pour tenir compte de sa direction et de son sens, on peut exprimer la tension du filin sous la forme : $\vec{T} = -T \frac{\vec{AM}}{l}$. On projette alors le principe fondamental de la dynamique $m\vec{a} = m\vec{g} + \vec{T} - 2m\vec{\Omega} \wedge \vec{v}$ dans la base cartésiennes proposée dans l'énoncé :

$$\begin{cases} m\ddot{x} = -T \frac{x}{l} - 2m\Omega (\cos \lambda \dot{z} - \sin \lambda \dot{y}) \\ m\ddot{y} = -T \frac{y}{l} - 2m\Omega \sin \lambda \dot{x} \\ m\ddot{z} = -mg - T \frac{z-l}{l} + 2m\Omega \dot{x} \cos \lambda \end{cases}$$

- D'après les calculs numériques de la 1^{ère} question, on a : $z \ll d$, $\dot{z} \ll \dot{d}$ et $\ddot{z} \ll \ddot{d}$ en tolérant un écart maximal de 10 % (on rappelle que d est la distance de la masse à l'axe, soit $d = \sqrt{x^2 + y^2}$). Comme de plus on cherche un terme correctif, on peut donc faire les approximations suivantes : $z \ll l$, $\dot{z} \ll \dot{y}$ et $\ddot{z} \ll \ddot{g}$.

D'autre part, on a également $\ddot{d} \ll g$ donc on pourra négliger le terme $2m\Omega \dot{x} \cos \lambda$ devant mg . Dans ces conditions, le système précédent devient :

$$\begin{cases} m\ddot{x} = -T \frac{x}{l} + 2m\Omega \sin \lambda \dot{y} \\ m\ddot{y} = -T \frac{y}{l} - 2m\Omega \sin \lambda \dot{x} \\ 0 = -mg + T \end{cases}$$

soit en posant $\omega_0^2 = \frac{T}{ml} = \frac{g}{l}$:

$$\begin{cases} \ddot{x} - 2\Omega \sin \lambda \dot{y} + \omega_0^2 x = 0 \\ \ddot{y} + 2\Omega \sin \lambda \dot{x} + \omega_0^2 y = 0 \\ T = mg \end{cases}$$

- $\underline{Z} = x + iy$ donc $\dot{\underline{Z}} = \dot{x} + i\dot{y}$ et $\ddot{\underline{Z}} = \ddot{x} + i\ddot{y}$. soit $\ddot{\underline{Z}} + 2i\Omega \sin \lambda \dot{\underline{Z}} + \omega_0^2 \underline{Z} = 0$.

8. L'équation caractéristique associée est : $r^2 + 2i\Omega \sin \lambda r + \omega_0^2 = 0$ dont les solutions sont :
 $r = -i \left(\Omega \sin \lambda \pm \sqrt{\Omega^2 \sin^2 \lambda + \omega_0^2} \right) \simeq -i (\Omega \sin \lambda \pm \omega_0)$ compte tenu des valeurs respectives de $\omega_0 = 0,386 \text{ rad.s}^{-1}$ et $\Omega = 7,27 \cdot 10^{-5} \text{ rad.s}^{-1}$ ($\omega_0 \gg \Omega$). Compte tenu des conditions initiales, on en déduit : $x = l \sin \alpha \cos \omega_0 t \sin \Omega \sin \lambda t$ et $y = l \sin \alpha \cos \omega_0 t \cos \Omega \sin \lambda t$.
9.
$$\underline{Z} = e^{-i\Omega \sin \lambda t} \left(l \sin \alpha \cos \omega_0 t + i l \sin \alpha \frac{\Omega \sin \lambda}{\omega_0} \cos \omega_0 t \right) = \underline{Z}' e^{-i\Omega \sin \lambda t}$$

 Dans le référentiel tournant à la vitesse angulaire $\Omega \sin \lambda$ autour de l'axe Oz , l'affixe du point est $\underline{Z}' = x' + iy'$ avec $x' = l \sin \alpha \cos \omega_0 t$ et $y' = l \sin \alpha \frac{\Omega \sin \lambda}{\omega_0} \cos \omega_0 t$. On a donc un mouvement elliptique, de demi grand axe $a = l \sin \alpha$, de demi petit axe $b = l \sin \alpha \frac{\Omega \sin \lambda}{\omega_0}$ donc très allongé puisque $\frac{\Omega \sin \lambda}{\omega_0} \ll 1$. Le pendule oscille à la pulsation ω_0 dans un plan qui tourne à la vitesse angulaire $\Omega \sin \lambda$ autour de l'axe vertical. Autrement dit, le plan d'oscillations du pendule tourne au cours du temps.
10. La période de rotation du plan d'oscillations vaut donc $T' = \frac{2\pi}{\Omega \sin \lambda}$. A la latitude de $48^\circ 51'$, on obtient : $T' = 115000 \text{ s} = 31 \text{ h } 52 \text{ min}$. On obtient pratiquement la même valeur que celle observée expérimentalement par Foucault. L'écart peut s'expliquer par exemple par une différence de la valeur de l'angle α et par d'autres problèmes liés au dispositif expérimental.
11. A l'équateur où $\lambda = 0$, T' tend vers l'infini, ce qui correspond au fait que le plan d'oscillations ne tourne pas, conformément à ce qui a été dit plus haut.
 Aux pôles, $\lambda = \pm 90^\circ$ donc $T' = 86400 \text{ s} = 24 \text{ h}$.
 A la latitude de 45° , $T' = 122000 \text{ s} = 33 \text{ h } 56 \text{ min}$.
12. Le résultat dépend de l'hémisphère dans lequel on se place : les sens de rotation du plan d'oscillations sont opposés. Il tourne vers l'est dans l'hémisphère nord et vers l'ouest dans l'hémisphère sud (à cause de $\sin \lambda$).

B.5 1. a) Dans le référentiel $\mathcal{R}_C$ galiléen, la Terre subit l'attraction gravitationnelle de l'astre A :

$$M_T \vec{a}_{T/\mathcal{R}_C} = -k M_A M_T \frac{\vec{AT}}{AT^3} \Rightarrow \vec{a}_{T/\mathcal{R}_C} = -k M_A \frac{\vec{AT}}{AT^3} = \vec{G}_A(T)$$

b) Dans le référentiel $\mathcal{R}_T$, le point P subit les forces suivantes :

- l'attraction gravitationnelle de la Terre,
- l'attraction gravitationnelle de l'astre A ,
- la force d'inertie d'entraînement due à la translation de $\mathcal{R}_G$ par rapport à $\mathcal{R}_C$,
- la force d'inertie d'entraînement due à la rotation de $\mathcal{R}_T$ par rapport à $\mathcal{R}_G$,
- les forces de contact $\vec{f}$,
- pas de force de coriolis car P est immobile dans $\mathcal{R}_T$.

La relation fondamentale de la dynamique dans $\mathcal{R}_T$ est :

$$m \vec{a}_{P/\mathcal{R}_T} = m \vec{G}_T(P) + m \vec{G}_A(P) - m \vec{a}_{T/\mathcal{R}_C} + m \Omega^2 \vec{HM} + \vec{f}$$

où le point H est le projeté orthogonal de P sur l'axe de rotation de la Terre.

c) Le terme $m \vec{X}(P)$ est $m \vec{G}_T(P) + m \Omega^2 \vec{HM}$: il s'agit du poids $m \vec{g}$. Le terme $\vec{C}_{AT}(P)$ est :

$$m \vec{C}_{AT}(P) = m \vec{G}_A(P) - m \vec{a}_{T/\mathcal{R}_C} = m(\vec{G}_A(P) - \vec{G}_A(T))$$

d'après l'expression de l'accélération de T .

2. a) Le terme $\vec{C}_{AT}(P)$ s'écrit :

$$\vec{C}_{AT}(P) = -kM_A \left(\frac{\vec{AP}}{AP^3} - \frac{\vec{AT}}{AT^3} \right)$$

On exprime d'abord les deux vecteurs :

$$\vec{AT} = -TA\vec{u}_x \text{ et } \vec{AP} = (x - TA)\vec{u}_x + \gamma\vec{u}_y + z\vec{u}_z$$

Le développement limité au premier ordre du terme $1/AP^3$ s'écrit :

$$AP^{-3} = ((x - TA)^2 + \gamma^2 + z^2)^{-3/2} = \frac{1}{TA^3} \left(1 - 2\frac{x}{TA} + \frac{x^2}{TA^2} + \frac{\gamma^2}{TA^2} + \frac{z^2}{TA^2} \right)^{-3/2}$$

soit finalement :

$$\frac{1}{PA^3} \simeq \frac{1}{TA^3} \left(1 + \frac{3x}{TA} \right)$$

le vecteur $\vec{C}_{AT}(P)$ s'écrit alors :

$$-\frac{kM_A}{TA^3} \left((x - TA) \left(1 + \frac{3x}{TA} \right) \vec{u}_x + \gamma \left(1 + \frac{3x}{TA} \right) \vec{u}_y + z \left(1 + \frac{3x}{TA} \right) \vec{u}_z + TA\vec{u}_x \right)$$

En développant et en ne gardant que les termes du premier ordre, on obtient l'expression de l'énoncé.

- b) Le terme dépendant de x , γ et z dans $\vec{C}_{AT}(P)$ a comme ordre de grandeur R_T/TA donc :

$$\frac{\|\vec{C}_{AT}\|}{\|\vec{G}_T\|} \simeq \frac{kM_A R_T}{TA^3} \frac{R_T^2}{kM_T} \simeq \frac{M_A}{M_T} \frac{R_T^3}{TA^3} \simeq 6.10^{-6}$$

Les forces sont représentées aux quatre points demandés sur la figure (S32.4)

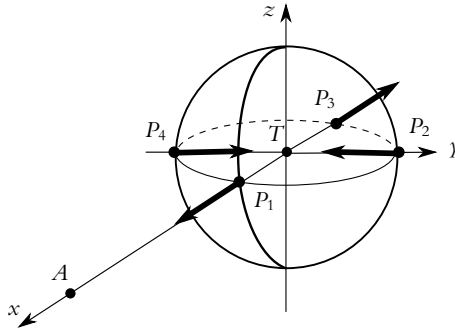


Figure S32.4

- c) On sait que $R_T^2 = x^2 + \gamma^2 + z^2$, donc $-2x^2 + \gamma^2 + z^2$ est égal à $R_T^2 - 3x^2$. Or avec la définition de Z_A , on peut écrire $x = R_T \cos Z_A$, d'où l'expression demandée.
3. a) Comme l'altitude z est petite devant R_T , on suppose que l'expression précédente de l'énergie potentielle est valable et on en déduit l'énergie potentielle gravitationnelle terrestre :

$$E_p = \frac{-kM_T m}{R_T + z} + kM_A m \frac{R_T^2}{2TA^3} (1 - 3 \cos^2 Z_A)$$

On effectue un développement limité à l'ordre 1 :

$$E_p = -km \left(\frac{M_T}{R_T} \left(1 - \frac{z}{R_T} \right) + \frac{M_A R_T^2}{2TA^3} (1 - 3 \cos^2 Z_A) \right)$$

soit :

$$\frac{E_p}{km} = z \frac{M_T}{R_T^2} - \frac{M_T}{R_T} + \frac{M_A R_T^2}{2TA^3} (1 - 3 \cos^2 Z_A)$$

- b) Pour l'équipotentielle E_{p0} , on déduit de l'expression précédente l'expression de z :

$$z = \frac{E_{p0}R_T^2}{kmM_T} + R_T + \frac{M_A R_T^4}{2M_T T A^3} (3 \cos^2 Z_A - 1)$$

Le maximum $z_{\max}$ est obtenu pour $Z_A = 0$ ou $Z_A = \pi$ et le minimum pour $Z_A = \pi/2$ soit :

$$z_{\max} = \frac{E_{p0}R_T^2}{kmM_T} + R_T + 2\frac{M_A R_T^4}{2M_T T A^3} \quad \text{et} \quad z_{\min} = \frac{E_{p0}R_T^2}{kmM_T} + R_T - \frac{M_A R_T^4}{2M_T T A^3}$$

Par différence, on obtient alors :

$$\Delta z = z_{\max} - z_{\min} = \frac{3M_A R_T^4}{2M_T T A^3}$$

L'application numérique donne $\Delta z = 0,56$ m pour la Lune et $\Delta z = 0,25$ m pour le soleil (effet deux fois moins important). On peut s'attendre à deux marées par jour en raison du $\cos^2 Z_A$.

Partie VI – Thermodynamique

Chapitre 33

- A.1** 1. Le manomètre mesure l'écart p_i entre la pression des pneumatiques P_i et la pression extérieure qu'on notera P_0 et qui vaut 1,0 atm.

2. En considérant le gaz à l'intérieur comme parfait et le volume constant, on a $\frac{P_1}{T_1} = \frac{nR}{V_1} = \frac{nR}{V_2} = \frac{P_2}{T_2}$ avec $P_1 = P_0 + p_1$ et $P_2 = P_0 + p_2$. On en déduit $p_2 = \frac{T_2}{T_1} (P_0 + p_1) - P_0$ soit 2,5 atm.

3. $\Delta p = 0,50$ atm soit 20 % qui est supérieur à l'écart maximal toléré dont il faut changer la pression en fonction de la saison.

- A.2** 1. L'échelle centésimale est définie par

$$\theta(x) = 100 \frac{V - V_0}{V_{100} - V_0} = \frac{a - bt + ct^2}{a - 100b + 10^4 c} t.$$

2. L'écart demandé s'exprime par $\Delta = \frac{b(100 - t) + c(t^2 - 10^4)}{a - 100b + 10^4 c}$.

3. On doit résoudre l'équation $\frac{d\Delta}{dt} = 0$ soit $3ct^2 - 2bt + 100b - 10^4 c = 0$ dont le discriminant vaut $4b^2 - 12c(100b - 10^4 c) = (2b - 300c)^2 + 3 \cdot 10^4 c^2 > 0$. On a donc deux solutions réelles $\frac{2b \pm \sqrt{(2b - 300c)^2 + 3 \cdot 10^4 c^2}}{6c}$.

- A.3** 1. On note T la température en F et t celle en °C. On cherche les coefficients a et b tels que $t = aT + b$, ce qui revient compte tenu des données à résoudre le système formé par $32a + b = 0$ et $212a + b = 100$. La solution est $t = \frac{5}{9} (T - 32) = 0,556 (T - 32)$.

2. On inverse la relation précédente et on obtient $T = 1,8t + 32$.

3. On applique la relation de la première question soit $t = 233^\circ\text{C}$.

4. On cherche la solution de $t = 1,8t + 32$ par exemple et on obtient $t = -40$.

- A.4** 1. D'après l'équation d'état : $P_C = \frac{RT_C}{V_C - b} - \frac{a}{V_C^2}$ (1).

La dérivée partielle de P par rapport à V est nulle, ce qui se traduit par $\frac{RT_C}{(V_C - b)^2} = \frac{2a}{V_C^3}$ (2).

De même pour la dérivée partielle seconde $\frac{2RT_C}{(V_C - b)^3} = \frac{6a}{V_C^4}$ (3).

En divisant la relation (2) par la (3) $\frac{V_C - b}{2} = \frac{V_C}{3}$ soit $b = \frac{V_C}{3}$.

En reportant dans (2), on en déduit $\frac{RT_C}{V_C - b} = \frac{4a}{3V_C^2}$ et dans (1) $a = 3P_C V_C^2$ puis dans (2) $R = \frac{8P_C V_C}{3T_C}$.

- En remplaçant a , b et R par les expressions obtenues, on a le résultat.
- On obtient une même expression pour tous les gaz indépendamment du point critique, ce qui correspond à la loi des états correspondants.

B.1 On utilisera un indice B (respectivement H) pour les variables relatives au gaz du bas (respectivement au gaz du haut).

- Les forces qui s'exercent sur le piston sont : son poids $m\vec{g}$, la force de pression exercée par le gaz du bas $\vec{F}_{p,B}$ dirigée vers le haut, et la force de pression exercée par le gaz du haut $\vec{F}_{p,H}$ dirigée vers le bas. On projette le principe fondamental de la dynamique à l'équilibre sur un axe vertical ascendant : $-mg + P_B S - P_H S = 0$, où S est la section du tube. La masse par unité de surface σ étant $\frac{m}{S}$, on obtient la relation cherchée :

$$P_B = P_H + \sigma g$$

L'application numérique donne :

$$P_H = 10 \text{ cmHg} = \frac{10}{76} 1,013 = 0,133 \text{ bar}$$

$$P_B = \left(\frac{10}{76} 1,013 \cdot 10^5 + 13600 \right) \text{ Pa} = 0,269 \text{ bar}$$

On a utilisé le fait que $76 \text{ cmHg} = 1 \text{ atm} = 1,013 \text{ bar}$.

- On cherche à déterminer la pression finale dans le compartiment du haut, ce qui permettra de déterminer le volume final de ce compartiment et donc le déplacement du piston. Le volume total, noté V , se conserve.

Pour le système gaz du haut (nombre de moles n_H) :

Etat initial	$P_H = P_1$ $V_H = \frac{V}{2}$ $T_H = 273 \text{ K}$	Etat final	P'_H V'_H $T'_H = 373 \text{ K}$
--------------	---	------------	--

Pour le système gaz du bas (nombre de moles n_B) :

Etat initial	$P_B = P_1 + \sigma g$ $V_B = \frac{V}{2}$ $T_B = 273 \text{ K}$	Etat final	P'_B V'_B $T'_B = 373 \text{ K}$
--------------	--	------------	--

L'équilibre du piston entraîne : $P'_B = P'_H + \sigma g$.

L'équation du gaz parfait donne, pour le gaz H :

$$n_H R = \frac{P_1 V_H}{T_1} = \frac{P'_H V'_H}{T_2} \quad \text{donc} \quad V'_H = \frac{T_2}{T_1} P_1 V_H$$

et pour le gaz B :

$$n_B R = \frac{(P_1 + \sigma g) V_B}{T_1} = \frac{(P'_H + \sigma g) V'_B}{T_2} \quad \text{donc} \quad V'_B = \frac{T_2}{T_1} \frac{P_1 + \sigma g}{P'_H + \sigma g} V_B$$

La conservation du volume entraîne : $V = V_H + V_B = V'_H + V'_B$ soit $2V_H = V'_H + V'_B$. On remplace V'_H et V'_B par les expressions obtenues plus haut :

$$2V_H = \frac{T_2 P_1}{T_1 P'_H} V_H + \frac{T_2}{T_1} \frac{P_1 + \sigma g}{P'_H + \sigma g} V_H$$

La pression P'_H vérifie l'équation du second degré :

$$2 \frac{T_1}{T_2} P_H^2 + P'_H \left(2 \frac{T_1}{T_2} \sigma g - 2P_1 - \sigma g \right) - P_1 \sigma g = 0$$

La seule solution positive est $P'_H = 1,99.10^4$ Pa.

Pour calculer le rapport $\frac{h'}{h}$ des hauteurs du piston dans l'état final et dans l'état initial, il suffit d'écrire :

$$\frac{h'}{h} = \frac{V'_H}{V_H} = \frac{P_1 T_2}{P'_H T_1} = 0,9$$

Le piston monte de donc de $0.1h$.

- B.2** 1. a) Initialement le piston est en haut. Lorsqu'il commence sa descente, la soupape s se ferme. Le système considéré est donc le gaz dans le cylindre de pression P , jusqu'à l'ouverture de s' pour V'_1 et $P = P_0$.

Etat initial $\left \begin{array}{l} P = P_a \\ V = V_M \\ T \\ n_C \end{array} \right $	$\left \right $	Etat final $\left \begin{array}{l} P = P_0 \\ V = V'_1 \\ T \\ n_C \end{array} \right $
---	------------------	--

On applique l'équation du gaz parfait pour trouver V'_1 :

$$V'_1 = \frac{P_a V_M}{P_0}$$

- b) La soupape s' étant ouverte, le système est constitué du gaz dans le cylindre et du gaz dans le réservoir. L'état final est obtenu lorsque le piston est au fond du cylindre.

Etat initial $\left \begin{array}{l} P_0 \\ V'_1 + V \\ T \\ n_C + n_R \end{array} \right $	$\left \right $	Etat final $\left \begin{array}{l} P_1 \\ V_m + V \\ T \\ n_C + n_R \end{array} \right $
--	------------------	---

On applique de nouveau l'équation du gaz parfait :

$$(n_C + n_R)RT = P_1(V + V_m) = P_0(V + V'_1)$$

En utilisant la relation $P_0 V'_1 = P_a V_M$, on obtient :

$$P_1 = P_a \frac{V_M}{V + V_m} + P_0 \frac{V}{V + V_m} \quad (1)$$

- c) Pour que la soupape s' s'ouvre, il faut que le volume V'_1 soit atteint avant que le piston ne soit complètement descendu, donc que : $V'_1 \geq V_m$. On en déduit : $P_0 \leq P_{\max}$ avec $P_{\max} = P_a \frac{V_M}{V_m}$.
2. Pour les allers et retours suivants le raisonnement est semblable. Pour obtenir la pression P_2 , il suffit de remplacer P_0 par P_1 , dans l'expression (1) :

$$P_2 = P_a \frac{V_M}{V + V_m} + P_1 \frac{V}{V + V_m}$$

On remplace P_1 par son expression :

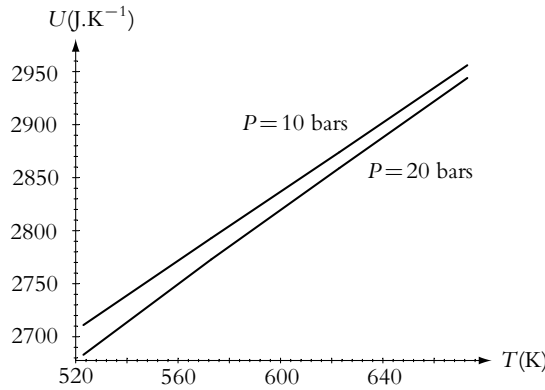
$$P_2 = P_a \left(\frac{V_M}{V + V_m} + \left(\frac{V_M}{V + V_m} \right)^2 \right) + P_0 \left(\frac{V}{V + V_m} \right)^2$$

On voit apparaître la formule de récurrence suivante :

$$P_n = P_a \sum_{i=1}^n \left(\frac{V_M}{V + V_m} \right)^i + P_0 \left(\frac{V}{V + V_m} \right)^n = P_a \frac{1 - \left(\frac{V_M}{V + V_m} \right)^n}{1 - \left(\frac{V_M}{V + V_m} \right)} + P_0 \left(\frac{V}{V + V_m} \right)^n$$

3. La limite pour n très grand donne : $\lim_{n \rightarrow \infty} P_n = P_a \frac{1}{1 - \frac{V_M}{V + V_m}} = P_{\max}$.
4. Applications numériques : $P_1 = 1,05$ bar et $P_{\max} = 25$ bar.

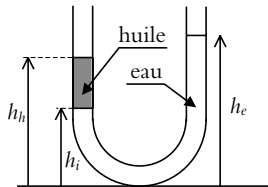
B.3 1. On obtient :



2. La variation d'énergie interne est proportionnelle à celle de la température du fait qu'on a une droite mais les deux courbes ne sont pas confondues donc l'énergie interne dépend de la pression. Le gaz n'est pas parfait : il faut tenir compte des interactions entre molécules.
3. Le calcul des pentes donne 28 J.mol^{-1} pour $P = 10 \text{ bars}$ et 29 J.mol^{-1} pour $P = 20 \text{ bars}$. Ce sont les valeurs des capacités thermiques. Dans le cas d'un gaz parfait monoatomique, on aurait $\frac{3}{2}R = 12,5 \text{ J.mol}^{-1}$. On a donc une pente plus élevée du fait des vibrations et des rotations.

Chapitre 34

A.1 Il faut établir trois relations entre h_h , h_i et h_e . La première est immédiate : $h_h - h_i = 2 \text{ cm}$. La seconde est obtenue en écrivant l'égalité des pressions dans l'eau et dans l'huile à l'interface, soit : $P_0 + \rho_{\text{huile}}(h_h - h_i)g = P_0 + \rho_{\text{eau}}(h_e - h_i)g$, ou encore (en utilisant la première relation) : $h_e - h_i = 2 \frac{\rho_{\text{huile}}}{\rho_{\text{eau}}}$ en centimètres.



La troisième est obtenue en écrivant que si le niveau de l'eau descend de x dans une branche, il monte de la même quantité dans l'autre, soit $h_h = H - x$ et $h_e = H + x$. Finalement, on trouve $x = \frac{\rho_{\text{huile}}}{\rho_{\text{eau}}}$ où x est exprimé en centimètres. On a donc $x = 0,6 \text{ cm}$ d'où $h_e = 10,6 \text{ cm}$, $h_i = 9,4 \text{ cm}$ et $h_h = 11,4 \text{ cm}$.

A.2 1. Il a été vu dans le cours qu'à une profondeur de 10 m la pression est le double de celle à la surface. En profondeur, les forces de pression qui s'exercent sur le corps et donc sur les poumons du plongeur sont plus importantes que s'il était à l'air libre. A l'équilibre, on peut supposer que la pression de l'air dans les poumons est égale à la pression extérieure. S'il remplit trop ses poumons d'air provenant de ses bouteilles et qu'il remonte rapidement, la pression extérieure exercée sur son corps risque de ne plus être suffisante pour s'opposer à la pression de l'air dans ses poumons et ces derniers risquent de se déchirer.

2. Pour calculer la pression à 10 m de profondeur, on applique la loi barométrique :

$$P = P_0 + \rho gh = 10^5 + 10^3 \cdot 9,8 \cdot 10 = 1,98 \cdot 10^5 \text{ Pa}$$

En supposant que le plongeur n'expire pas d'air pendant la remontée, le poumon contient n moles d'air. Si on note V le volume à la profondeur h et V_0 à la surface :

$$nRT = PV = P_0 V_0 \Rightarrow V_0 = \frac{PV}{P_0} = 1,98V = 5,94 \text{ L}$$

Le poumon devrait doubler de volume, ce qui est impossible.

- A.3** 1. D'après l'énoncé, la loi de température est $T(z) = T_0 - az$ avec $T(0) = T_0 = 293 \text{ K}$ et :

$$a = -\frac{T(8850) - T(0)}{8850} = 6,8 \cdot 10^{-3} \text{ K.m}^{-1}$$

2. L'équation de la statique des fluides s'écrit $dP = -\rho g dz$ avec $\rho = \frac{m}{V} = \frac{MP}{RT}$. On en déduit l'équation différentielle $\frac{dP}{P} = -\frac{Mg}{R} \frac{dz}{T_0 - az}$ en tenant compte de la dépendance avec l'altitude de

la température. On intègre entre 0 où la pression est P_0 et z et on obtient $P = P_0 \left(1 - \frac{az}{T_0}\right)^{\frac{Mg}{aR}}$.

L'application numérique au sommet de l'Everest donne :

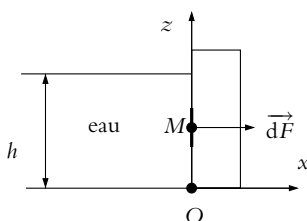
$$P = 0,316 P_0$$

3. On différentie la température en fonction de l'altitude $dT = -adz$ soit $dz = -\frac{dT}{a}$. On reporte alors dans l'équation de la statique des fluides, ce qui donne : $\frac{dP}{P} = \frac{Mg}{Ra} \frac{dT}{T}$. Cette équation s'intègre en : $PT^{-\frac{Mg}{aR}} = \text{cste}$. Grâce à l'équation d'état du gaz parfait, on obtient : $PV^k = \text{cste}$ où $k = \frac{Mg}{Mg - Ra}$. L'application numérique donne $k = 1,25$.

- A.4** La force élémentaire $\vec{dF}$ exercée par l'eau sur la surface dS (en trait gras sur la figure) est :

$$\vec{dF} = P(M) \vec{dS} = P(M) dS \vec{e}_x$$

où $\vec{e}_x$ est le vecteur unitaire de l'axe Ox et $dS = \ell dz$.



Il faut donc calculer la pression en M avec la loi barométrique :

$$P(M) = P_0 + \rho gz \quad \text{soit :} \quad dF = (P_0 + \rho gz) \ell dz$$

La force totale est obtenue en intégrant :

$$F = \int_0^h (P_0 + \rho gz) \ell dz = P_0 h \ell + \frac{1}{2} \rho g h^2 \ell$$

L'application numérique donne : $F = 2,49 \cdot 10^6 \text{ N}$. Attention à prendre la pression P_0 en Pascal.

- A.5** 1. On applique la loi barométrique :

$$P(z) = -\rho gz + P_0$$

donc à une profondeur $z = -10$ km et avec $P_0 = 1$ bar, la pression est :

$$P = 1,01 \cdot 10^8 \text{ Pa} = 1,01 \cdot 10^3 \text{ bar}$$

2. a) On considère une masse m constante de fluide homogène $m = \rho V$, d'où

$$\left(\frac{\partial V}{\partial P} \right)_T = \left(\frac{\partial \left(\frac{m}{\rho} \right)}{\partial P} \right)_T = \frac{-m}{\rho^2} \left(\frac{\partial \rho}{\partial P} \right)_T = \frac{-V}{\rho} \left(\frac{\partial \rho}{\partial P} \right)_T$$

d'où :

$$\chi_T = \frac{-1}{V} \left(\frac{\partial V}{\partial P} \right)_T = \frac{1}{\rho} \left(\frac{\partial \rho}{\partial P} \right)_T$$

- b) La loi fondamentale de la statique s'écrit :

$$\frac{dP}{dz} = -\rho g \quad (1)$$

or :

$$\frac{dP}{dz} = \left(\frac{\partial P}{\partial \rho} \right)_T \frac{\partial \rho}{\partial z}$$

Si on considère l'eau en isotherme, l'équation (1) s'écrit :

$$\frac{d\rho}{dz} = -\rho^2 g \chi_T \Rightarrow \frac{d\rho}{\rho^2} = -g \chi_T dz \Rightarrow \frac{1}{\rho} = g \chi_T z + \frac{1}{\rho_0}$$

3. a) On obtient alors l'équation donnant la pression :

$$\frac{dP}{dz} = -\rho g = -\frac{\rho_0 g}{1 + \chi_T \rho_0 g z} \Rightarrow P = P_0 - \frac{1}{\chi_T} \ln(1 + \chi_T \rho_0 g z)$$

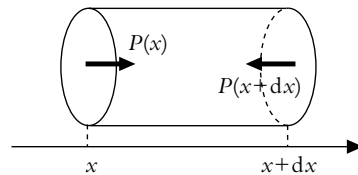
- b) Application numérique : $P = 1,03 \cdot 10^8$ Pa soit une différence de 2% par rapport au modèle de l'eau incompressible

B.1 1. a) Les forces qui s'exercent sur cette petite tranche d'air de masse dm dans le référentiel tournant $\mathcal{R}'$ lié au tube sont :

- le poids ;
- la force d'inertie d'entraînement ou force centrifuge ;
- les forces de pression exercée par le reste du gaz.

En effet, $\mathcal{R}'$ n'est pas galiléen et il faut tenir compte des forces d'inertie. Cependant, la force de Coriolis est nulle puisque l'air est à l'équilibre dans le référentiel $\mathcal{R}'$.

- b) On néglige le poids, donc il ne reste que la force centrifuge et les forces de pression. On considère le volume mésoscopique $d\tau_M$ d'air situé entre les plans d'abscisses x et $x + dx$: il a une épaisseur dx et une surface S égale à celle du tube.



Le volume $d\tau_M$ est donc égal à $S dx$ et la masse dm à $\rho d\tau_M$. La force inertie d'entraînement s'exerçant sur ce volume est donc : $d\vec{f}_{ie} = dm \omega^2 x \vec{u}_x$, $\vec{u}_x$ étant le vecteur unitaire de l'axe Ox . La relation fondamentale de la dynamique à l'équilibre appliquée à $d\tau_M$, en projection sur Ox , s'écrit :

$$SP(x) - SP(x + dx) + dm \omega^2 x = 0$$

Un développement au premier ordre en dx , donne :

$$-S \frac{dP}{dx} dx + \rho S dx \omega^2 x = 0 \quad \text{soit :} \quad \frac{dP}{dx} = \rho S \omega^2 x$$

- c) On intègre la relation précédente : $P(x) = \frac{1}{2} \rho \omega^2 x^2 + \text{cte}$. On détermine la constante par continuité de la pression en $x = L$, soit :

$$P(x) = P_a + \frac{1}{2} \rho \omega^2 (x^2 - L^2)$$

2. Comme on l'a signalé à la question précédente, on néglige la variation de pression sur la hauteur du tube dans le gaz : cette variation due au poids est négligeable dans un gaz sur des petites hauteurs (voir cours). On en déduit que la pression en N est obtenue en appliquant l'expression trouvée précédemment en $x = 0$, soit $P_N = P_a - \frac{1}{2}\rho\omega^2 L^2$. On applique alors la loi barométrique au liquide entre N et un point au niveau de la surface libre. La dénivellation est donnée par :

$$P_Q = P_a = P_N + \mu gh \Rightarrow h = \frac{1}{2\mu g}\rho\omega^2 L^2$$

- B.2** 1. En procédant comme à l'exercice précédent, on obtient, en projetant l'équation traduisant l'équilibre d'un petit volume de fluide sur les vecteurs de base des coordonnées cylindriques $\vec{u}_r$ et $\vec{u}_z$ (vertical ascendant) :

$$\begin{cases} \frac{\partial P}{\partial r} = \rho\omega^2 r \\ \frac{\partial P}{\partial z} = -\rho g \end{cases}$$

ou encore : $\vec{\text{grad}}P = \rho(\vec{g} - \omega^2 \overrightarrow{HM})$ où H est le projeté orthogonal de M sur l'axe de rotation.

La surface libre correspond à la surface sur laquelle la pression P est constante. Il faut donc que le second terme soit nul soit $\tan \alpha = \frac{g}{\omega^2 r} = \frac{dz}{dr}$. On en déduit l'équation $z = \frac{\omega^2}{2g}r^2 + C$ avec C une constante.

Si la notion de gradient n'a pas encore été vue, on peut simplement intégrer les équations précédentes : $P(r, z) = \frac{1}{2}\rho\omega^2 r^2 - \rho g z + \text{cste}$ (on vérifie que les deux dérivées partielles de cette fonction correspondent bien à celles que l'on attend). La surface libre correspond à $P(r, z) = P_0$, ce qui donne

$$\text{bien : } z = \frac{\omega^2}{2g}r^2 + C.$$

Pour déterminer la valeur de la constante, il suffit d'écrire la conservation de la quantité d'eau soit

$$\pi R^2 h = \int_0^R 2\pi r z dr \text{ avec } h = 10 \text{ cm. On en déduit } C = h - \frac{\omega^2 R^2}{4g}.$$

2. Le centre est découvert si la constante C est nulle soit d'après la valeur obtenue à la question précédente pour une vitesse de rotation $\omega = \sqrt{\frac{4gh}{R^2}} = 40 \text{ rad.s}^{-1}$.

3. On a débordement si $z(R) > h' = 40 \text{ cm}$ soit $\omega > \sqrt{\frac{4g(h' - h)}{R^2}} = 69 \text{ rad.s}^{-1}$.

4. $P = P_{\text{atm}} + \rho g h_c - \frac{\rho\omega^2 R^2}{4} = P_{\text{atm}} + 364 \text{ Pa}$.

- B.3** 1. a) On ne peut pas utiliser la poussée d'Archimède car la demi-sphère est posée au fond d'un récipient. En appliquant la loi barométrique :

$$P(M) = P(O) - \rho g(z_M - z_0) = P(O) - \rho g a \cos \theta$$

- b) L'axe Oz étant un axe de symétrie, les composantes horizontales des forces de pression s'annulent. La force résultante est portée par Oz et c'est pour cela qu'on projette la force élémentaire sur l'axe Oz . La force de pression exercée par le liquide en M est $d\vec{F}_M = -P(M)d\vec{S}_M$, soit en projection sur Oz :

$$dF_z = -P(M) \cos \theta dS_M$$

- c) Avec l'expression de dS_M et de $P(M)$, on peut exprimer dF_z :

$$dF_z = -2\pi a^2 (P(O) - \rho g \cos \theta) \cos \theta \sin \theta d\theta$$

Pour calculer la force totale, il faut intégrer la force sur la demi-sphère :

$$F_z = -2\pi a^2 \int_{\theta=0}^{\theta=\pi/2} (P(O) - \rho g \cos \theta) \cos \theta \sin \theta d\theta = -2\pi a^2 \left(\frac{P(O)}{2} - \frac{\rho g}{3} \right)$$

2. a) On peut appliquer l'expression démontrée à la question précédente, sachant qu'alors :

$$P(O) = P_0 + \rho g(h + a) \quad \text{d'où :} \quad F_z = -2\pi a^2 \left(\frac{P_0}{2} + \frac{\rho g h}{2} + \frac{\rho g a}{6} \right) = -32 \text{ N}$$

- b) S'il n'y avait que la pression atmosphérique, la force serait $-\pi a^2 P_0$. On voit d'après l'expression précédente, que sous l'eau la force est plus importante (de 0,3 N).

- B.4** 1. La sphère de bois étant en équilibre, la résultante des forces qui s'exercent sur elle est nulle soit

$$M_b \vec{g} - \iint_{M \in S} P(M) \vec{dS} + \vec{R} = \vec{0}$$

en notant M_b la masse de la sphère et $\vec{R}$ la réaction du fond du bassin.

La difficulté est de calculer la résultante des forces de pression. On ne peut en effet pas appliquer le théorème d'Archimède car la pression n'est pas continue à la surface de la sphère. Il faut donc calculer explicitement l'intégrale.

L'air est à une pression uniforme donc il ne contribue pas à l'intégrale :

$$\iint_{M \in S} P_a \vec{dS} = P_a \iint_{M \in S} \vec{dS} = \vec{0}$$

La pression en un point de cote z dans l'eau est donnée par : $P(z) = P_a + \rho_e g(H - z)$. En fonction de l'angle θ des coordonnées sphériques, on obtient :

$$P(\theta) = P_a + \rho_e g(H + R(\cos \theta_0 - \cos \theta))$$

où θ_0 représente la valeur de θ au niveau de l'ouverture circulaire : $\sin \theta_0 = \frac{r}{R}$ et

$$\cos \theta_0 = -\sqrt{1 - \frac{r^2}{R^2}}.$$

Les symétries du problème imposent à la résultante de n'avoir qu'une composante F_z suivant $\vec{u}_z$. On projette sur $\vec{u}_z$: $dF_z = (P(\theta) - P_a) \times (-\cos \theta) \times dS$. On peut découper la sphère en couronnes comprises entre θ et $\theta + d\theta$ de surface : $dS = 2\pi R^2 \sin \theta d\theta$. Il vient :

$$F_z = -\rho_e g 2\pi R^2 \int_0^{\theta_0} (H + R(\cos \theta_0 - \cos \theta)) \cos \theta \sin \theta d\theta$$

On trouve :

$$F_z = \rho_e g 2\pi R^2 \left(\frac{H}{2} (\cos^2 \theta_0 - 1) + R \left(\frac{1}{3} - \frac{1}{2} \cos \theta_0 + \frac{1}{6} \cos^3 \theta_0 \right) \right)$$

La force exercée par la sphère sur le fond du bassin est donc

$$\vec{R} = M_b \vec{g} \left(1 - \frac{\rho_e}{\rho} \left(\frac{3H}{4} (\cos^2 \theta_0 - 1) + R \left(\frac{1}{2} - \frac{3}{4} \cos \theta_0 + \frac{1}{6} \cos^3 \theta_0 \right) \right) \right) \vec{u}_z$$

2. L'application numérique donne $\|\vec{R}\| = 170 \text{ N}$.
3. La sphère monte à la surface avant d'émerger lorsque la résultante précédente s'annule donc dès que H atteint la valeur :

$$H_c = \frac{2R}{3 \sin^2 \theta_0} \left(1 - \frac{2\rho_b}{\rho_e} - \frac{3}{2} \cos \theta_0 + \frac{1}{2} \cos^3 \theta_0 \right)$$

En remplaçant $\cos \theta_0$ par son expression en fonction de r et R , on obtient :

$$H_c = \frac{2R^3}{3r^2} \left(\left(1 + 2\frac{\rho_b}{\rho_e} \right) + \left(2 + \frac{r^2}{R^2} \right) \sqrt{1 - \frac{r^2}{R^2}} \right)$$

- B.5** 1. a) Si $k = 0$, la transformation est isotherme.

- b) *Homogène* signifie que les variables d'état et en particulier la masse volumique sont les mêmes en tout point. *Isotrope* signifie que les propriétés sont les mêmes quelle que soit la direction. Dans le cas de l'atmosphère, il y a une direction particulière qui est la verticale et la masse volumique dépend de l'altitude. Donc l'atmosphère n'est ni homogène, ni isotrope.

- c) Si on appelle m la masse d'un volume mésoscopique et V son volume, $V = mv(z) = nM_{\text{air}}v(z)$. L'équation d'état du gaz parfait est donc :

$$p(z)V = nRT \Rightarrow p(z)v(z) = \frac{RT}{M_{\text{air}}}$$

- d) Sachant que $v(z) = \frac{1}{\rho(z)}$ où $\rho(z)$ est la masse volumique, la loi de la statique des fluides devient :

$$\frac{dp}{dz} = -\rho(z)g = \frac{-g}{v(z)}$$

- e) On veut calculer $\frac{dT}{dz}$. La loi (34.20) donne : $T = Cp^k$ où $C = p_0^{-k}T_0$, soit :

$$\frac{dT}{dz} = Ckp^{k-1} \frac{dp}{dz}$$

Avec $p^{k-1} = \frac{T}{Cp}$, cette équation s'écrit :

$$\frac{dT}{dz} = k \frac{T}{p} \frac{dp}{dz} = k \frac{M_{\text{air}}v(z)}{R} \left(\frac{-g}{v(z)} \right) = -\frac{M_{\text{air}}kg}{R} \Leftrightarrow \delta = \frac{M_{\text{air}}kg}{R} = 5,1 \cdot 10^{-3} \text{ K.m}^{-1}$$

- f) On en déduit : $T(z) = T_0 - \delta z$.

- g) On applique la loi des gaz parfaits au niveau du sol et à l'altitude z : $p(z)V(z) = nRT(z)$ et $p_0V_0 = nRT_0$. On effectue le rapport de ces deux équations en utilisant la relation (34.20) :

$$\frac{V(z)}{V_0} = \frac{T(z)}{T_0} \frac{p(z)}{p_0} \Rightarrow \frac{V(z)}{V_0} = \left(\frac{T(z)}{T_0} \right)^{1-\frac{1}{k}}$$

d'où avec l'expression de $T(z)$:

$$\frac{V(z)}{V_0} = \left(\frac{T_0 - \delta z}{T_0} \right)^{1-\frac{1}{k}} \quad (2)$$

2. a) La force F est égale à la poussée d'Archimède moins le poids du dihydrogène. Etant donné que l'on suppose les deux gaz parfaits, le gaz dans le ballon et l'air déplacé occupant le même volume correspondent au même nombre de moles n_0 .

$$F = n_0(M_{\text{air}} - M_{\text{H}_2})g$$

- b) Pour que le ballon décolle, il faut : $F > m_B g$ soit $n_0 > n_{\text{min}}$ avec :

$$n_{\text{min}} = \frac{m_B}{M_{\text{air}} - M_{\text{H}_2}}$$

- c) L'altitude maximale est atteinte lorsque le ballon aura atteint son volume maximal $10V_0$. En utilisant l'expression (2), on en déduit l'altitude h correspondante :

$$h = \frac{T_0}{\delta} \left(1 - \left(\frac{10V_0}{V_0} \right)^{\frac{k}{k-1}} \right) = 19,2 \text{ km}$$

Cette altitude est au-delà de la limite supérieure de la troposphère donc le modèle utilisé n'est plus valable.

B.6

1. a) Les seules forces qui agissent sur un petit volume de fluide à l'équilibre sont son poids et les forces de pression. Le poids est vertical donc la résultante des forces de pression aussi. La pression ne dépend ni de x ni de y .

- b) L'équilibre du petit volume de fluide sur Oz donne $dP = -\rho g dz$ avec $\rho = \frac{n'M}{V} = \frac{MP}{RT}$ en utilisant la loi des gaz parfaits. Par intégration suivant z , on en déduit $P = P_0 e^{-\frac{Mgz}{RT}}$.

- c) Par définition, on a $n = \frac{n_{\text{mol}}}{V} = \frac{n'N_A}{V}$ en notant N_A le nombre d'Avogadro. En reportant dans l'équation d'état $PV = n'RT$ et en utilisant la relation $R = N_A k$, on obtient le résultat.

- d) D'après la question précédente, on a $n = \frac{P}{kT}$, ce qui permet d'obtenir le résultat en notant $n_0 = \frac{P_0}{kT}$ et $E(z) = \frac{Mgz}{N_A}$.

e) $E(z)$ est l'énergie potentielle d'une molécule et kT l'énergie d'agitation thermique.

f) L'application numérique de la question 2. donne $\frac{P(z)}{P_0} = 0,30$.

g) Le rapport cherché s'exprime par $r = \frac{\int_0^{1000} n(z) dz}{\int_0^{+\infty} n(z) dz} = 1 - e^{-\frac{1000Mg}{RT}}$ dont l'application numérique donne $r = 0,70$.

2. a) Le rapport vaut par définition $\gamma = \frac{C_P}{C_V} = 1,4$.

b) On formule l'hypothèse que le produit PV^γ est constant et on a l'équation des gaz parfaits $PV = nRT$. On en déduit que $V = \frac{nRT}{P}$ et que le produit $P^{1-\gamma}T^\gamma$ est constant.

c) La différentielle logarithmique de la relation précédente donne $\frac{dT}{T} = \frac{\gamma-1}{\gamma} \frac{dP}{P}$

d) En reportant la relation de la question précédente dans l'équation de la statique des fluides, on en déduit $\frac{dP}{P} = -\frac{Mg}{RT} dz = \frac{\gamma}{\gamma-1} \frac{dT}{T}$ soit $\left(\frac{dT}{dz}\right)_{\text{adia}} = \frac{Mg(1-\gamma)}{R\gamma}$. On a donc un gradient constant de la température avec l'altitude. Il est donc possible de donner sa valeur soit $-9,78 \cdot 10^{-3} \text{ K.m}^{-1}$.

e) La lecture de la température à 10 km donne -45° C . Le gradient obtenu par la pente de la courbe est donc ici $\left(\frac{dT}{dz}\right)_{\text{réel}} = -5,7 \cdot 10^{-3} \text{ K.m}^{-1}$. Cette valeur est supérieure à la valeur obtenue dans l'hypothèse adiabatique.

f) La valeur de q est obtenue par une relation analogue à celle utilisée avec γ soit $\left(\frac{dT}{dz}\right)_{\text{réel}} = \frac{Mg(1-q)}{Rq}$. On en déduit $q = \frac{Mg}{Mg + R\left(\frac{dT}{dz}\right)_{\text{réel}}} = 1,20$.

g) La valeur $\frac{dT}{dz}$ est constante donc par intégration, on en déduit $T = T_0 - 5,7 \cdot 10^{-3} z$.

h) Par une démonstration analogue à celle de question 2, on en déduit

$$P = P_0 \left(\frac{T}{T_0}\right)^{\frac{q}{q-1}}$$

i) En reportant l'expression de $T(z)$ dans la relation de la question précédente, on a

$$P = P_0 \left(1 - \frac{5,7 \cdot 10^{-3}}{T_0} z\right)^{\frac{q}{q-1}}$$

j) L'application numérique avec $T_0 = 288 \text{ K}$ à $z = 10 \text{ km}$ donne une température de $T = 231 \text{ K}$ et une pression de $P = 0,27 \text{ atm}$.

Chapitre 35

A.1 1. La vitesse de libération est obtenue pour une énergie mécanique nulle :

$$E_m = \frac{1}{2}mv_l^2 - \frac{GMm}{R} = 0 \Rightarrow v_l = 2\sqrt{\frac{MG}{R}}$$

M étant la masse de la planète et R son rayon.

On rappelle que la vitesse quadratique moyenne est $u = \sqrt{\frac{3k_B T}{m}}$, m étant la masse d'une particule.

Planète	v_l en km.s^{-1}
Mercuré	4,24
Vénus	10,4
Terre	11,2
Mars	5,01

Les vitesses quadratiques moyennes sont : $u_{H_2} = 1,94 \text{ km.s}^{-1}$ et $u_{N_2} = 0,52 \text{ km.s}^{-1}$. Il faut se souvenir que la vitesse quadratique moyenne est une moyenne, donc de nombreuses particules sont plus rapides. D'après les données précédentes, la vitesse quadratique est beaucoup plus proche des vitesses de libération pour Mercure et Mars que pour la Terre et Vénus. Les particules s'échappent plus facilement de l'attraction de Mars et Mercure, ce qui peut expliquer qu'il n'y ait pas d'atmosphère.

2. Pour que N_2 échappe à l'attraction terrestre, il faut que $u_{N_2} = v_l$, soit $T = 1,4 \cdot 10^5 \text{ K}$.

A.2 1. Les particules de V_1 passant dans V_2 entre t et $t + dt$ sont celles contenues dans un cylindre de base s et de hauteur $v dt \vec{u}_x$ compte tenu des hypothèses très simplifiées considérées. Son volume est $vs dt$ et le nombre de ces particules $N_1 \frac{vs dt}{6V}$, le facteur 6 est lié au fait que seulement un sixième des particules ont une vitesse $v \vec{u}_x$. On en déduit $dN_{1 \rightarrow 2} = \frac{N_1 vs}{6V} dt$.

2. De même, on a $dN_{2 \rightarrow 1} = \frac{N_2 vs}{6V} dt$. On en déduit

$$dN_1 = -dN_2 = -dN_{1 \rightarrow 2} + dN_{2 \rightarrow 1} = \frac{vs}{6V} (N_2 - N_1) dt$$

3. Comme $N_1 + N_2 = N_a$, on obtient $dN_1 = \frac{vs}{6V} (N_a - 2N_1) dt$. On en déduit l'équation différentielle $\frac{dN_1}{dt} + \frac{vs}{3V} N_1 = \frac{vs}{6V} N_a$.

La résolution de cette équation différentielle du premier ordre à coefficients constants donne $N_1 = \frac{N_a}{2} + C \exp\left(-\frac{vst}{3V}\right)$. Or à $t = 0, 0 \text{ s}$, $N_1 = N_a = \frac{N_a}{2} + C$ donc $C = \frac{N_a}{2}$. Finalement $N_1 = \frac{N_a}{2} \left(1 + \exp\left(-\frac{t}{\tau}\right)\right)$ avec $\tau = \frac{3V}{vs}$ et $N_2 = \frac{N_a}{2} \left(1 - \exp\left(-\frac{t}{\tau}\right)\right)$.

4. Avec $v = \sqrt{\frac{3RT}{M}}$, on en déduit $\tau = \frac{V}{s} \sqrt{\frac{3M}{RT}}$ donc le temps caractéristique est proportionnel à $\sqrt{M}$.

5. La diffusion étant proportionnelle à $\sqrt{M}$ donc si on a deux types de particules de masse différente, on peut enrichir les mélanges en bloquant le système avant la fin.

B.1 1. La paroi et le piston sont fixes donc leur énergie cinétique est nulle. L'énergie cinétique totale est donc uniquement celle de la particule, elle est conservée au cours du choc : la norme de la vitesse de la particule est inchangée.

2. Avant le choc, la particule a une vitesse $\vec{v}_i = v \vec{u}_x$ et après le choc $\vec{v}_f = -v \vec{u}_x$, donc la variation de la composante selon Ox de la quantité de mouvement est :

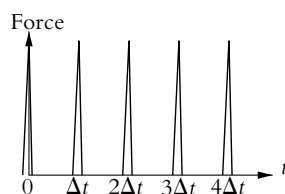
$$\Delta p_x \vec{u}_x = \mu(\vec{v}_f - \vec{v}_i) = -2\mu v \vec{u}_x$$

L'intervalle entre deux chocs successifs sur le piston est l'intervalle mis pour parcourir $2L$, soit $\Delta t = \frac{2L}{v}$.

3. La particule exerce une force sur le piston à des intervalles de temps réguliers Δt . En supposant que le premier choc a lieu à $t = 0$, on obtient la figure ci-contre.

4. D'après la figure, pendant $\tau = n\Delta t$, le piston subit $n + 1$ chocs, donc la variation de quantité de mouvement totale du piston est :

$$\Delta p'_x(\tau) = -(n + 1)\Delta p_x = 2(n + 1)\mu v$$



D'où la force moyenne :

$$F_m = \lim_{n \rightarrow \infty} \frac{2(n+1)\mu\nu}{n\Delta t} = \frac{2\mu\nu}{\Delta t} = \frac{\mu\nu^2}{L}$$

5. Pour obtenir la force moyenne exercée par toutes les particules il suffit de sommer les $F_{m,i}$, soit :

$$F_m = \sum_{i=1}^N F_{m,i} = \sum_{i=1}^N \frac{\mu v_i^2}{L} \quad \text{d'où :} \quad F_m = \frac{2N}{L} \overline{E_c}$$

6. En utilisant l'expression précédente de $\overline{E_c}$ et celle donnée par l'énoncé, on trouve l'expression de F_m :

$$F_m = \frac{Nk_B T}{L}$$

or la pression exercée sur le piston est $P = \frac{F_m}{S}$ d'où :

$$PSL = Nk_B T \Rightarrow PV = Nk_B T$$

On retrouve l'équation d'état du gaz parfait.

Chapitre 37

A.1 1. Faux. Il faudrait que le gaz soit parfait.

2. Faux. Il faut en plus que l'état initial et l'état final soient des états d'équilibre avec l'extérieur.

3. Vrai. Cf. cours.

4. Faux. Suivant les conditions initiales, le gaz peut se refroidir ou se réchauffer.

5. Faux. Par exemple, au cours d'une détente isotherme d'un gaz parfait, le gaz fournit du travail, reçoit du transfert thermique et sa température ne varie pas. Autre exemple : si on veut faire fondre de manière isobare et isotherme un glaçon, il faut lui fournir un transfert thermique.

6. Faux. D'après le premier principe $\Delta U = 0$ mais U n'est pas fonction que de la température en général, donc T n'est pas forcément constante. Pour un gaz parfait en revanche, le résultat est vrai.

7. a) Faux. Adiabatique ne veut pas dire isotherme.

b) Vrai car $W < 0$ et $Q = 0$.

8. a) Vrai. Le travail à fournir est égal à l'aire sous la courbe représentative de la transformation dans le diagramme de Watt (P, V). Or la pente de l'adiabatique étant supérieure à celle de l'isotherme à l'état initial, il faut fournir plus de travail pour comprimer le système de manière adiabatique.

b) Faux. $W = \Delta U$ uniquement si la transformation est adiabatique.

A.2 Le système considéré est constitué du calorimètre et des deux masses d'eau qu'on y met.

1. En négligeant la capacité thermique du calorimètre, le bilan énergétique se traduit par : $\Delta H_{\text{masse 1}} + \Delta H_{\text{masse 2}} = 0$ puisqu'il n'y a pas de travail ni d'échange avec l'extérieur. On obtient l'équation $m_1 c_0 (\theta_f - \theta_1) + m_2 c_0 (\theta_f - \theta_2) = 0$ soit $\theta_f = \frac{m_2 \theta_2 + m_1 \theta_1}{m_1 + m_2} = 32,8^\circ \text{C}$.

2. L'écart de température entre la valeur théorique et la valeur constatée prouve qu'on ne peut pas négliger la capacité thermique du calorimètre.

D'autre part, la température constatée étant inférieure à la température théorique, on peut en déduire qu'on commence par mettre l'eau la plus froide dans le calorimètre.

3. En notant m_0 la valeur en eau du récipient, le nouveau bilan énergétique s'écrit

$$(m_0 + m_1) c_0 (\theta'_f - \theta_1) + m_2 c_0 (\theta'_f - \theta_2) = 0 \quad \text{soit} \quad m_0 = m_2 \frac{\theta_2 - \theta'_f}{\theta'_f - \theta_1} - m_1 = 22,5 \text{ g.}$$

A.3 Le système considéré est constitué du calorimètre, de l'eau qu'il contient et de la barre métallique qu'on y introduit.

Comme il n'y a ni transfert thermique avec l'extérieur ni travail, le bilan énergétique se traduit par $\Delta H = 0 \Leftrightarrow (\mu + m_e) c_e (\theta_f - \theta_e) + m_m c_m (\theta_f - \theta_m) = 0$.

On en déduit $c_m = c_e \frac{(\mu + m_e) (\theta_f - \theta_e)}{m_m (\theta_m - \theta_f)} = 982 \text{ J.K}^{-1}.\text{kg}^{-1}$.

A.4 1. La transformation étant quasistatique, le travail des forces de pression s'écrit $\delta W = -PdV$ (la pression du système étant infiniment proche d'un état d'équilibre est pratiquement égale à la pression extérieure). De plus, on a PV^k constant donc par différentiation logarithmique $\frac{dP}{P} + k \frac{dV}{V} = 0$ soit

$VdP = -kPdV$. Comme $d(PV) = PdV + VdP = (1 - k)PdV$, on en déduit $\delta W = \frac{d(PV)}{k-1}$ et

par intégration $W = \frac{P_1 V_1 - P_0 V_0}{k-1}$.

2. Comme on a un gaz parfait, on peut écrire $dU = C_V dT$. Par ailleurs, la relation de Mayer s'écrit $C_P - C_V = nR$ soit avec $\gamma = \frac{C_P}{C_V}$, on en déduit $C_V = \frac{nR}{\gamma-1}$ et $dU = \frac{nRdT}{\gamma-1}$. Par intégration et en utilisant la loi des gaz parfaits, on a $\Delta U = \frac{nR}{\gamma-1} (T_1 - T_0) = \frac{P_1 V_1 - P_0 V_0}{\gamma-1}$.

Alors en appliquant le premier principe $\Delta U = W + Q$, on en déduit le transfert thermique $Q = \left(\frac{1}{\gamma-1} - \frac{1}{k-1} \right) (P_1 V_1 - P_0 V_0) = \frac{k-\gamma}{(k-1)(\gamma-1)} (P_1 V_1 - P_0 V_0)$ qu'on peut écrire sous la forme demandée en utilisant la loi des gaz parfaits avec $C = \frac{(k-\gamma)nR}{(\gamma-1)(k-1)}$.

3. C est homogène à une capacité thermique.

4. a) Pour $k = \gamma$, $C = 0$, ce qui signifie qu'il n'y a pas d'échange thermique : la transformation est adiabatique.

b) Pour $k = 0$, $C = \frac{\gamma nR}{\gamma-1} = C_P$ et la transformation est isobare.

c) Pour k infini, $C = \frac{nR}{\gamma-1} = C_V$ et la transformation est isochore.

d) Pour $k = 1$, C tend vers l'infini. Le système peut emmagasiner tout transfert thermique sans variation de température, la transformation est donc isotherme.

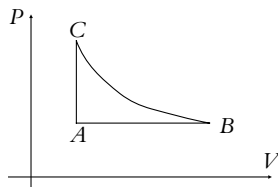
5. Pour $k = \gamma$, on trouve $Q = 0$ donc $W = \frac{P_1 V_1 - P_0 V_0}{\gamma-1}$.

Pour $k = 0$, on obtient $W = P_0 V_0 - P_1 V_1 = nR(T_0 - T_1)$.

Pour k infini, $W = 0$.

Pour $k = 1$, retrouver le résultat n'est pas simple.

A.5 1. Le cycle a l'allure suivante :



2. a est homogène au produit d'une pression par un volume au carré donc s'exprime en Pa.m^6 dans les unités du système international.

b est homogène à un volume donc s'exprime en m^3 dans les unités du système international.

3. Pour le point B, on connaît le volume $V_B = 5,0 \text{ L}$ et la température $T_B = 1000 \text{ K}$. Avec l'équation d'état, on en déduit la pression $P_B = \frac{RT_B}{V_B - b} - \frac{a}{V_B^2} = 1,7 \cdot 10^6 \text{ Pa}$ (ou 17 bars).

Pour le point C, on connaît le volume $V_C = 1,0 \text{ L}$ et la température $T_C = 1000 \text{ K}$. Avec l'équation d'état, on en déduit la pression $P_C = \frac{RT_C}{V_C - b} - \frac{a}{V_C^2} = 8,4 \cdot 10^6 \text{ Pa}$ (ou 84 bars).

Enfin pour le point A, on connaît le volume $V_A = 1,0 \text{ L}$ et la pression $P_A = 8,4 \cdot 10^6 \text{ Pa}$. Avec l'équation d'état, on en déduit la pression $T_A = \frac{1}{R} \left(P_A + \frac{a}{V_A^2} \right) (V_A - b) = 210 \text{ K}$.

4. Les transformations étant quasistatiques, on peut identifier la pression extérieure à la pression du gaz et écrire le travail élémentaire $\delta W = -PdV$.

Pour AB qui est isobare, on a $W_{AB} = -P_A (V_B - V_A) = -6,67 \text{ kJ}$.

Pour BC qui est isotherme, on a $W_{BC} = - \int_{V_B}^{V_C} \frac{RT_B}{V - b} - \frac{a}{V^2} dV$ soit en explicitant l'intégration

$$W_{BC} = -RT_B \ln \frac{V_C - b}{V_B - b} - a \left(\frac{1}{V_C} - \frac{1}{V_B} \right) = 13,48 \text{ kJ}.$$

Pour CA qui est isochore, le travail est nul $W_{CA} = 0,0 \text{ kJ}$.

5. Au cours d'un cycle, le premier principe appliqué au gaz s'écrit $\Delta U = Q + W = 0$. Or ici $W = W_{AB} + W_{BC} + W_{CA} = 6,81 \text{ kJ}$ donc $Q = -6,81 \text{ kJ}$.

6. On calcule la variation d'énergie interne du gaz en utilisant l'expression de l'énergie interne fournie puis on déduit le transfert thermique à partir de l'expression du premier principe $\Delta U_i = W_i + Q_i$ soit $Q_i = \Delta U_i - W_i$.

Pour la transformation AB, on a $\Delta U_{AB} = \frac{3}{2}R(T_B - T_A) - a \left(\frac{1}{V_B} - \frac{1}{V_A} \right) = 9,96 \text{ kJ}$. On en déduit $Q_{AB} = \Delta U_{AB} - W_{AB} = 16,63 \text{ kJ}$.

Pour la transformation BC, on a $\Delta U_{BC} = \frac{3}{2}R(T_C - T_B) - a \left(\frac{1}{V_C} - \frac{1}{V_B} \right) = -0,11 \text{ kJ}$. On en déduit $Q_{BC} = \Delta U_{BC} - W_{BC} = -13,59 \text{ kJ}$.

Pour la transformation CA, on a $\Delta U_{CA} = \frac{3}{2}R(T_A - T_C) - a \left(\frac{1}{V_A} - \frac{1}{V_C} \right) = -9,85 \text{ kJ}$. On en déduit $Q_{CA} = \Delta U_{CA} - W_{CA} = -9,85 \text{ kJ}$.

7. On peut vérifier que $Q = Q_{AB} + Q_{BC} + Q_{CA}$.

A.6 Le système considéré est le gaz.

1. On décompose la première transformation en deux étapes :

- état initial : P_0, T_0 ;
- état intermédiaire : $P_1 = 20P_0, T_0$;
- état final : P_0, T_1 .

La seconde étape est adiabatique quasi-statique donc on peut utiliser la loi de Laplace :

$$P_1^{1-\gamma} T_0^\gamma = P_0^{1-\gamma} T_1^\gamma \Rightarrow T_1 = T_0 \left(\frac{P_1}{P_0} \right)^{\frac{1-\gamma}{\gamma}}$$

L'application numérique donne : $T_1 = 115 \text{ K}$.

2. On reprend le raisonnement précédent pour T_2 et T_1 , ce qui donne :

$$T_2 = T_1 \left(\frac{P_1}{P_0} \right)^{\frac{1-\gamma}{\gamma}} = T_0 \left(\frac{P_1}{P_0} \right)^{2 \frac{1-\gamma}{\gamma}}$$

De même :

$$T_n = T_{n-1} \left(\frac{P_1}{P_0} \right)^{\frac{1-\gamma}{\gamma}} = T_0 \left(\frac{P_1}{P_0} \right)^{n \frac{1-\gamma}{\gamma}}$$

L'application numérique donne, pour $n = 5$: $T_5 = 3,78 \text{ K}$. C'est une méthode qui peut sembler très efficace pour refroidir un gaz mais en fait on a supposé la transformation réversible et le gaz parfait. A très basse température, si le système est encore gazeux, il faut utiliser une équation d'état plus réaliste telle que l'équation de Van Der Waals.

3. On décompose la $n^{\text{ème}}$ transformation en deux étapes

- état initial : P_0, T_{n-1} ;
- état intermédiaire : $P_1 = 20P_0, T_{n-1}$;
- état final : P_1, T_n .

On note n_m le nombre de moles du gaz. La variation d'énergie interne ΔU_1 est nulle pour la transformation isotherme d'après la 1^{ère} loi de Joule. Pour la seconde transformation :

$$\Delta U_2 = \frac{n_m R}{\gamma - 1} (T_n - T_{n-1}) = \frac{n_m R}{\gamma - 1} T_0 \left(\frac{P_1}{P_0} \right)^{(n-1) \frac{1-\gamma}{\gamma}} \left(\left(\frac{P_1}{P_0} \right)^{\frac{1-\gamma}{\gamma}} - 1 \right)$$

Finalement, la variation d'énergie interne est : $\Delta U = \Delta U_2$ puisque $\Delta U_1 = 0$.

Pour la première transformation qui est quasistatique isotherme, le travail est :

$$W_1 = n_m R T_n \ln \left(\frac{P_1}{P_0} \right) = n_m R T_0 \left(\frac{P_1}{P_0} \right)^{\frac{1-\gamma}{\gamma}} \ln \left(\frac{P_1}{P_0} \right)$$

et avec le premier principe :

$$Q_1 = \Delta U_1 - W_1 = -W_1$$

Pour la seconde transformation qui est adiabatique $Q_2 = 0$, donc $W_2 = \Delta U_2$. D'où finalement : $W = W_1 + W_2$ et $Q = Q_1$.

Application numérique : $\Delta U = -106 \text{ J}$, $W = -12, 2 \text{ J}$ et $Q = -94, 1 \text{ J}$.

B.1 1. Dans l'état final, la température de l'enceinte A est T_1 puisqu'on a chauffé cette enceinte et la température de l'enceinte B toujours en contact avec le thermostat est T_0 . Par ailleurs, l'équilibre mécanique se traduit par l'égalité des pressions dans les deux enceintes soit avec l'équation des gaz parfaits $P_f = \frac{RT_1}{V_A} = \frac{RT_0}{V_B}$. On a conservation du volume c'est-à-dire que les volumes V_A et V_B vérifient $V_A + V_B = 2V_0$. De l'égalité des pressions, on tire $V_B = V_A \frac{T_0}{T_1}$ donc en reportant dans la conservation du volume, on obtient $V_A = \frac{2V_0 T_1}{T_1 + T_0}$ puis $V_B = \frac{2V_0 T_0}{T_1 + T_0}$. Finalement la pression finale vaut $P_f = \frac{R(T_0 + T_1)}{2V_0}$.

2. Comme les gaz sont parfaits, les variations d'énergie interne ne dépendent que des variations de température. Par conséquent, dans l'enceinte B , la variation d'énergie interne est nulle $\Delta U_B = 0$. Pour l'enceinte A , on a $\Delta U_A = C_V (T_1 - T_0)$. De la relation de Mayer $C_P - C_V = nR$ et de la définition de γ rappelée dans l'énoncé, on obtient $C_V = \frac{nR}{\gamma - 1}$ et finalement $\Delta U_A = \frac{R}{\gamma - 1} (T_1 - T_0)$ (ici $n = 1$ mole). La variation d'énergie interne de l'ensemble s'obtient en appliquant l'extensivité de l'énergie interne soit $\Delta U = \Delta U_A + \Delta U_B = \frac{R}{\gamma - 1} (T_1 - T_0)$.

3. L'évolution de l'enceinte B est quasistatique et à température constante : elle est donc isotherme. On en déduit l'expression du travail reçu par l'enceinte B :

$$W = - \int_{V_0}^{V_B} P dV = \int_{V_0}^{V_B} RT_0 \frac{dV}{V} = -RT_0 \ln \frac{V_B}{V_0} = RT_0 \ln \frac{T_1 + T_0}{2T_0}.$$

On déduit le transfert thermique Q_1 reçu par l'enceinte B en provenance du thermostat (il n'y a pas d'échange avec l'enceinte A qui est calorifugée) en appliquant le premier principe $\Delta U_B = 0 = Q_1 + W$ soit $Q_1 = RT_0 \ln \frac{2T_0}{T_1 + T_0}$.

4. On considère que le gaz de l'enceinte A et la résistance constitue le système calorifugé. L'application du premier principe au gaz de l'enceinte A donne $\Delta U_A = -W + Q_2$ où Q_2 est le transfert thermique fourni par la résistance au gaz de l'enceinte A soit $Q_2 = \frac{R}{\gamma - 1} (T_1 - T_0) + RT_0 \ln \frac{T_1 + T_0}{2T_0}$.

B.2 1. Le système (S) considéré est constitué du gaz, du piston, du cylindre et de la masse posée sur le piston.

A l'état initial, le gaz est à la pression P_0 , température T_0 , volume $V_0 = hs$. A l'état final, le gaz est à la pression P_1 , température T_1 , volume $V_1 = h_1s$. Le piston ayant une masse m_0 et n'étant pas surmonté par un gaz, $P_0 = \frac{m_0g}{s}$ et $P_1 = \frac{3m_0g}{s}$ (on rappelle que pour déterminer les pressions, il suffit d'écrire l'équilibre mécanique du piston).

La transformation est adiabatique réversible, donc on peut appliquer la loi de Laplace au gaz parfait contenu dans le cylindre :

$$P_0 V_0^\gamma = P_1 V_1^\gamma \Leftrightarrow h_1 = h \left(\frac{1}{3} \right)^{\frac{1}{\gamma}} = 4,6 \text{ cm}$$

Pour obtenir la température, on applique aussi une des lois de Laplace :

$$P_0^{1-\gamma} T_0^\gamma = P_1^{1-\gamma} T_1^\gamma \Leftrightarrow T_1 = T_0 3^{\frac{\gamma-1}{\gamma}} = 374 \text{ K}$$

2. On étudie le même système. L'état initial est le même que précédemment et dans l'état final, le gaz est à la pression P_2 , son volume est V_2 et sa température T_2 . Puisque la masse totale déposée est la même, la pression finale P_2 est égale à P_1 . D'autre part, la masse est posée dès le début de l'évolution, donc la pression extérieure exercée sur le gaz est constante et égale à P_1 .

On applique le premier principe au système (S) :

$$\Delta U_{\text{gaz}} + \Delta U_{\text{enceinte}} + \Delta U_{\text{piston}} + \Delta E_{\text{cpiston}} = W + Q$$

Le piston est immobile au départ et à l'arrivée et on suppose que les capacités thermiques du piston et de l'enceinte sont négligeables. La transformation est adiabatique donc $Q = 0$. La pression extérieure est constante donc $W = P_1(V_0 - V_2)$. On obtient alors :

$$\Delta U_{\text{gaz}} = \frac{nR}{\gamma - 1} (T_2 - T_0) = P_1(V_0 - V_2)$$

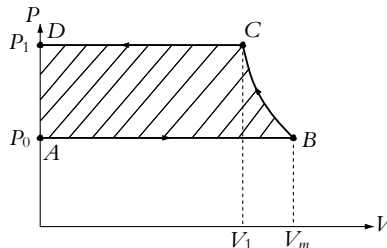
Or $nRT_0 = P_0 V_0$ et $nRT_2 = P_1 V_2$, d'où

$$T_2 = \frac{3(\gamma - 1) + 1}{\gamma} T_0 = 429 \text{ K}$$

et $h_2 = 5,2 \text{ cm}$.

B.3 1. a) Le système est le gaz qui entre dans le corps de la pompe pendant la première phase. L'évolution étudiée comporte les trois étapes suivantes :

- admission du gaz de A à B à la pression P_0 ;
- compression de B à C ;
- refoulement de C à D à la pression P_1 .



b) La pression extérieure est constante, donc le travail des forces de pression du gaz à droite du piston est :

$$W_{\text{droite}} = W_{\text{droite,AB}} + W_{\text{droite,BC}} + W_{\text{droite,CD}} = P_0(V_B - V_A + V_C - V_B + V_D - V_C) = 0$$

Effectivement, puisque le piston effectue un aller-retour, le volume balayé est le même à l'aller (travail résistant) et au retour (travail moteur).

La force exercée par le moteur pour déplacer le piston est $F = (P - P_0)s$, où s est la section du piston, donc le travail fourni par le moteur est :

$$W_{\text{moteur}} = - \int_{ABCD} (P - P_0) dV = - \int_{ABCD} P dV - W_{\text{droite}} = - \int_{ABCD} P dV$$

W_{moteur} est donc l'aire du cycle. C'est ce qu'on appelle habituellement le travail de l'opérateur.

2. a) On étudie le gaz dans le corps de pompe entre B et C. La loi polytropique et l'équation du gaz parfait donnent :

$$\begin{cases} P_B V_B^k = P_C V_C^k \\ P_B V_m = nRT_B \quad \text{et} \quad P_C V_C = nRT_C \end{cases}$$

En combinant ces expressions, on obtient :

$$\left(\frac{P_B}{P_C} \right) = \left(\frac{V_B}{V_C} \right)^k \Rightarrow \left(\frac{P_B}{P_C} \right)^{1-k} = \left(\frac{T_B}{T_C} \right)^{-k}$$

Comme on veut calculer k , on prend le logarithme de cette expression, ce qui donne :

$$k = \frac{\ln \frac{P_B}{P_C}}{-\ln \frac{T_B}{T_C} + \ln \frac{P_B}{P_C}} = \frac{\ln \frac{P_0}{P_1}}{-\ln \frac{T_0}{T_1} + \ln \frac{P_0}{P_1}} = 1,2$$

- b) On a vu dans précédemment que $W_{\text{moteur}} = \int (-P dv)$. On décompose le calcul suivant les trois transformations élémentaires :

- de A à B, $P = P_0$:

$$W_{\text{moteur},AB} = -P_0 \int_0^{V_m} dV = -P_0 V_m = -nRT_0$$

- de B à C :

$$W_{\text{moteur},BC} = \int_{V_m}^{V_1} (-P dV) = -P_0 V_m^k \int_{V_m}^{V_1} \frac{dV}{V^k} = \frac{P_0 V_m^k}{k-1} \left[V^{-k+1} \right]_{V_m}^{V_1}$$

ce qui peut s'écrire :

$$W_{\text{moteur},BC} = \frac{P_0 V_m}{k-1} \left(\left(\frac{V_m}{V_1} \right)^{k-1} - 1 \right)$$

Or :

$$\left(\frac{V_m}{V_1} \right)^{k-1} = \left(\frac{V_m}{V_1} \right)^k \frac{V_1}{V_m} = \frac{P_1 V_1}{P_0 V_m} = \frac{T_1}{T_0} \quad \text{d'où finalement : } W_{\text{moteur},BC} = \frac{P_0 V_m}{k-1} \left(\frac{T_1}{T_0} - 1 \right)$$

- de C à D, $P = P_1$ et :

$$W_{CD} = - \int_{V_1}^0 P_1 dV = P_1 V_1 = nRT_1$$

Ainsi, finalement :

$$W_{\text{moteur}} = \frac{k}{k-1} nR(T_1 - T_0)$$

La pompe aspire une mole d'air à chaque aller et retour. On prend donc $n = 1$ mol dans l'expression précédente et on trouve $W_{\text{moteur}} = 5,04 \text{ kJ} \cdot \text{mol}^{-1}$.

- c) Sur la figure ci-dessous, on a décomposé les étapes AB et CD de la transformation.

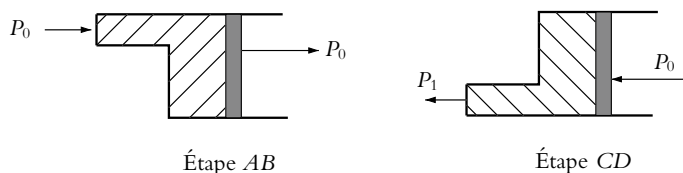


Figure S37.1 Étape AB : le volume balayé est V_m ; Étape CD : le volume balayé est V_1 .

On a montré que le travail de l'air à droite du piston W_{droite} est nul sur un aller et retour.

Il reste à calculer le travail de l'atmosphère à gauche de la tranche d'air aspiré pendant AB et le travail de l'air du réservoir à gauche de la tranche d'air transvasée pendant CD .

Dans le premier cas, le volume balayé est V_m , à la pression P_0 : le travail moteur est $W_{\text{gauche},AB} = P_0 V_m$. Pendant le refoulement, la pression du gaz qui s'oppose au refoulement dans le réservoir est P_1 , le volume balayé est V_1 et le travail résistant est donc $W_{\text{gauche},CD} = -P_1 V_1$.

On applique le premier principe au gaz admis dans le cylindre pendant un aller et retour :

$$\Delta U = U_D - U_A = W_{\text{total}} + Q = W_{\text{droite}} + W_{\text{gauche},AB} + W_{\text{gauche},CD} + W_{\text{moteur}} + Q$$

Or $U_A + P_0 V_m = H_A$, $U_D + P_1 V_1 = H_D$ et $W_{\text{droite}} = 0$ soit finalement :

$$\Delta H = W_{\text{moteur}} + Q$$

On en déduit :

$$Q = \frac{nR\gamma}{\gamma - 1}(T_1 - T_0) - W_{\text{moteur}} = -2,1 \cdot 10^3 \text{ kJ} \cdot \text{mol}^{-1}$$

Q est négatif, donc le système perd du transfert thermique.

3. La puissance du moteur est égale au travail du moteur pour une mole multiplié par le débit massique et divisé par la masse d'une mole :

$$\mathcal{P}_{\text{moteur}} = \frac{DW_{\text{moteur}}}{M} = 2,26 \text{ kW}$$

B.4 1. Le cycle est formé de deux isochores BC et DF et de deux adiabatiques réversibles CD et FB . Il est représenté sur la figure ci-contre.

2. Un tour correspond à un cycle donc, puisque le régime est de 7000 tours.min⁻¹, la durée d'un cycle est :

$$t_c = \frac{60}{7000} = 8,6 \text{ ms}$$

Pendant un cycle, le piston effectue un aller et retour, donc 2 courses. La distance parcourue est donc $d = 2 \times 39,2 \text{ mm} = 78,4 \text{ mm}$. On en déduit la vitesse moyenne du piston : $v_{\text{moy}} = \frac{d}{t_c} = 9,2 \text{ m} \cdot \text{s}^{-1}$.

La vitesse quadratique moyenne des molécules est : $u = \sqrt{\frac{3RT}{M}}$, soit en prenant $T = T_F$, on trouve pour l'air : $u \simeq 500 \text{ m} \cdot \text{s}^{-1}$.

Les valeurs numériques calculées montrent que $v_{\text{moy}} \ll u$. Le piston se déplace donc à très faible vitesse par rapport aux molécules : on peut considérer la transformation lente pour les molécules donc quasistatique.

3. La compression est adiabatique réversible, on peut donc utiliser la loi de Laplace :

$$P_B V_C^\gamma = P_F V_D^\gamma$$

d'où l'expression du taux de compression :

$$a = \frac{V_D}{V_C} = \left(\frac{P_B}{P_F} \right)^{\frac{1}{\gamma}} = 3,6$$

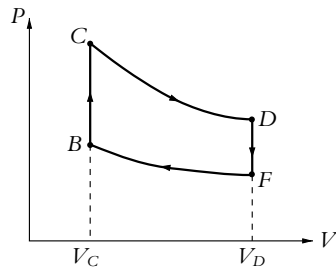
4. Le travail entre B et C est nul puisque il s'agit d'une transformation isochore. Il en est de même entre D et F . Le travail W_t reçu par le gaz au cours d'un cycle est donc : $W_t = W_{CD} + W_{FB}$.

Comme les transformations CD et FD sont adiabatiques, on peut écrire :

$$W_{CD} = \Delta U_{CD} = \frac{nR}{\gamma - 1}(T_D - T_C) \quad \text{et} \quad W_{FB} = \Delta U_{FB} = \frac{nR}{\gamma - 1}(T_B - T_F)$$

D'autre part, en utilisant la loi des gaz parfaits, $nR = \frac{P_F V_D}{T_F}$, on en déduit l'expression du travail fourni W par le moteur qui est l'opposé du travail reçu W_t :

$$W = -\frac{P_F V_D}{(\gamma - 1)T_F}(T_D - T_C + T_B - T_F)$$



5. La combustion BC étant isochore, on peut écrire :

$$Q_{BC} = \Delta U_{BC} = \frac{nR}{\gamma - 1}(T_C - T_B) = \frac{P_F V_D}{(\gamma - 1)T_F}(T_C - T_B)$$

6. Il s'agit d'un moteur, donc le travail cédé W est positif et la chaleur reçue Q_{BC} aussi. Le rendement s'écrit :

$$\eta = \frac{W}{Q_{BC}} = -\frac{T_D - T_C + T_B - T_F}{T_C - T_B} = 1 + \frac{T_F - T_D}{T_C - T_B}$$

Il faut exprimer les températures en fonction des volumes pour faire apparaître le taux de compression. On utilise la loi de Laplace pour les transformations CD et FB :

$$T_C V_C^{\gamma-1} = T_D V_D^{\gamma-1} \quad \text{et} \quad T_B V_C^{\gamma-1} = T_F V_D^{\gamma-1}$$

d'où :

$$T_D = T_C a^{1-\gamma} \quad \text{et} \quad T_F = T_B a^{1-\gamma}$$

En reportant dans l'expression de η , on trouve :

$$\eta = 1 + \frac{T_B a^{1-\gamma} - T_C a^{1-\gamma}}{T_C - T_B} = 1 - a^{1-\gamma}$$

7. A 45 km.h^{-1} , la puissance du moteur est $P = 4,4 \text{ kW}$ et un cycle est décrit en $t_c = 8,6 \text{ ms}$. Le travail cédé est : $W = P t_c = 37,8 \text{ J}$. On en déduit que la chaleur reçue lors de la combustion est : $Q_{BC} = \frac{W}{\eta} = 94,6 \text{ J}$.

8. Sur un cycle, le volume d'essence consommé est : $V_{\text{cycle}} = \frac{Q_{BC}}{q}$. Le moteur a un régime de $n_t = 7000 \text{ tours.min}^{-1}$, donc le volume d'essence consommé par minute est :

$$V_{\text{min}} = \frac{n_t Q_{BC}}{q} = 22 \text{ cm}^3$$

Pour parcourir 100 km à la vitesse de $45,5 \text{ km.h}^{-1}$, il faut $\Delta t = 132 \text{ min}$. La consommation pour 100 km est donc $132 \times V_{\text{min}} = 2,9 \text{ L}$, ce qui semble raisonnable.

B.5 1. Méthode de Rüchardt :

- a) La présence de la bille impose une force supplémentaire $m \vec{g}$ au système. En terme de pression c'est-à-dire de force ramenée à l'unité de surface, on a une surpression $\Delta P = \frac{mg}{S} = 814 \text{ Pa}$. On en déduit que la surpression est négligeable devant la pression atmosphérique puisque $\Delta P \ll P_0 = 1,010^5 \text{ Pa}$.

- b) Par la définition proposée, on a $\alpha = 3,333 \text{ V.bar}^{-1}$ donc $\Delta V = \alpha \frac{\Delta P}{P_0} = 27 \text{ mV}$. La précision étant de l'ordre du millivolt ne pose pas de soucis ici.

- c) On applique le théorème de Millman en v_- soit $v_- = \frac{\frac{E_0}{R_1} + \frac{s}{R_2}}{\frac{1}{R_1} + \frac{1}{R_2}} = \frac{R_2 E_0 + R_1 s}{R_1 + R_2}$. Par ailleurs,

on utilise la relation du pont diviseur de tension pour obtenir la tension v_+ soit $v_+ = \frac{u}{2}$. Or on a les DEUX hypothèses que l'amplificateur opérationnel est idéal ET qu'il fonctionne en régime linéaire donc $\epsilon = 0$. On en déduit avec les expressions précédentes $\frac{R_2 E_0 + R_1 s}{R_1 + R_2} = \frac{u}{2}$ donc $s = \frac{(R_1 + R_2)u - 2R_2 E_0}{2R_1}$.

- d) On a $u = \alpha P$ donc $s = aP + b$ avec $a = \frac{(R_1 + R_2)\alpha}{2R_1}$ et $b = -\frac{R_2 E_0}{R_1}$.

- e) On traduit les deux conditions de l'énoncé $s = 0$ pour $P = P_{\text{atm}}$ soit $\frac{(R_1 + R_2)\alpha}{2R_1} P_0 = \frac{R_2 E_0}{R_1}$ et $\frac{s}{\Delta P} = \frac{(R_1 + R_2)\alpha}{2R_1}$.

De la seconde, on tire $R_2 = \frac{2R_1s}{\alpha\Delta P} - R_1 = 179 \text{ k}\Omega$ et en reportant dans la première, on a $E_0 = \frac{\alpha(R_1 + R_2)P_0}{2R_2} = 1,68 \text{ V}$.

- f) Par symétrie (on se reportera à la partie sur l'électromagnétisme pour plus de détails), on en déduit que la force est dirigée suivant Oz . On projette les éléments de surface sphérique $r^2 \sin \theta d\theta d\varphi$ sur Oz .

On en déduit $F_z = 2\pi P_0 r^2 \int_0^{\frac{\pi}{2}} \sin \theta \cos \theta d\theta = \pi r^2 P_0 = sP_0$.

- g) Le système étudié est la bille dans le référentiel terrestre considéré comme galiléen. Les forces sont le poids, la force de pression atmosphérique et la force de pression intérieure. Le principe fondamental de la dynamique projeté sur la verticale orientée vers le haut donne $m\ddot{z} = -mg + s(P - P_0)$.
- h) Les relations de Laplace ne sont valables que pour un gaz parfait subissant une transformation adiabatique et quasistatique.

- i) Par différentiation de la relation de Laplace $PV^\gamma = \text{constante}$, on obtient $\frac{dP}{P} + \gamma \frac{dV}{V} = 0$ soit avec les approximations proposées $\frac{P - P_0}{P_0} + \gamma \frac{s\dot{z}}{V_0} = 0$. On en déduit $\Delta P = -\frac{\gamma s P_0 \dot{z}}{V_0}$.

- j) En reportant dans l'équation différentielle $\ddot{z} + \frac{\gamma P_0 s^2}{mV_0} z = -g$ et en utilisant le fait ΔP est proportionnel à $\dot{z}$, on en déduit $\ddot{P} + \frac{\gamma P_0 s^2}{mV_0} \left(P - P_0 - \frac{mg}{s} \right) = 0$.

- k) On obtient des oscillations dont la période est $T_0 = 2\pi \sqrt{\frac{mV_0}{\gamma P_0 s^2}}$. La position d'équilibre est donnée par la solution particulière soit $z_e = -\frac{mV_0 g}{\gamma P_0 s^2}$.

- l) De l'expression de la période T_0 , on déduit $\gamma = \frac{4\pi^2 mV_0}{P_0 s^2 T_0^2}$.

- m) L'application numérique donne $\gamma = 1,3$.

- n) L'application numérique pour la position d'équilibre donne $z_e = 0,31 \text{ m}$ soit une valeur un peu différente des observations. On en déduit que le modèle n'est pas entièrement satisfaisant.

2. Méthode de Clément-Desormes :

- a) Les coefficients thermoélastiques sont définis par $\alpha = \frac{1}{V} \left(\frac{\partial V}{\partial T} \right)_P$, $\beta = \frac{1}{P} \left(\frac{\partial P}{\partial T} \right)_V$ et

$$\chi_T = -\frac{1}{V} \left(\frac{\partial V}{\partial P} \right)_T.$$

- b) En reportant les expressions des dérivées partielles en fonction des coefficients thermoélastiques dans les expressions de l et h , on obtient $l = TP\beta$ et $h = -TV\alpha$.

- c) En différentiant la pression P en fonction de la température T et du volume V , on a $dP = \left(\frac{\partial P}{\partial T} \right)_V dT + \left(\frac{\partial P}{\partial V} \right)_T dV$. On reporte cette différentielle dans l'expression de δQ soit

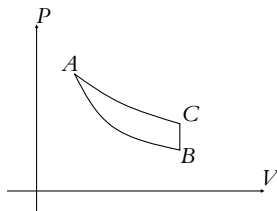
$$\delta Q = C_V dT + l dV = \left(C_P + h \left(\frac{\partial P}{\partial T} \right)_V \right) dT + h \left(\frac{\partial P}{\partial V} \right)_T dV \text{ ce qui permet d'obtenir par identification } C_P - C_V = -h \left(\frac{\partial P}{\partial T} \right)_V = TVP\alpha\beta.$$

- d) La relation des gaz parfaits $PV = nRT$ permet de calculer $\alpha = \beta = \frac{1}{T}$ et $\chi_T = \frac{1}{P}$. En reportant ces résultats dans l'expression de la question précédente, on en déduit la relation de Mayer $C_P - C_V = nR$. De même, avec les relations de la question b, on en déduit $l = P$ et $h = -V$.

- e) Si la transformation est isotherme, on a $dT = 0$ et $\delta Q = l dV = h dP$. On en déduit donc $\left(\frac{\partial P}{\partial V}\right)_T = \frac{l}{h}$.

Pour une transformation adiabatique, on a $\delta Q = 0$ donc $C_V dT = -l dV$ et $C_P dT = -h dP$. En faisant le rapport membre à membre de ces deux relations, on obtient $\left(\frac{\partial P}{\partial V}\right)_{\text{adia}} = \gamma \frac{l}{h} = \gamma \left(\frac{\partial P}{\partial V}\right)_T$. Ce résultat montre que la pente d'une adiabatique est, en valeur absolue, γ fois plus grande que celle d'une isotherme ($\gamma > 1$).

- f) Le système est le gaz après la fermeture du robinet pour pouvoir considérer un système fermé, ceci est nécessaire pour écrire les principes de la thermodynamique.
- g) La transformation AB est rapide, on peut considérer qu'il n'y a pas d'échange thermique donc que la transformation est adiabatique.
- h) La transformation BC est isochore car le robinet est fermé et il n'y a pas de variation de volume.
- i) On écrit le premier principe dans le cas où il n'y a pas de travail car la transformation est isochore soit $dU = \delta Q$. On utilise le fait que le gaz est parfait soit $dU = \frac{nR}{\gamma - 1} dT$. Par ailleurs, les hypothèses données permettent d'écrire $\delta Q = -k(T - T_0) dt$. En reportant ces deux expressions dans le premier principe, on obtient $\frac{nR}{\gamma - 1} dT = -k(T - T_0) dt$ soit $\frac{dT}{dt} + \frac{(\gamma - 1)k}{nR} (T - T_0) = 0$. C'est bien la forme de l'équation différentielle cherchée avec $\tau = \frac{nR}{(\gamma - 1)k}$.
- j) La transformation est isochore donc la pression P et la température T sont proportionnelles. On obtient la même équation différentielle $\frac{dP}{dt} + \frac{(\gamma - 1)k}{nR} (P - P_0) = 0$.
- k) La représentation du cycle est la suivante :



- l) La tension et la pression sont proportionnelles donc on en déduit l'expression du rapport $\gamma = \frac{P_B - P_A}{P_C - P_A} = \frac{u_B - u_A}{u_C - u_A} = 1,2$.

Chapitre 38

A.1 On peut considérer que la température du lac ne varie pas et dans l'état final, le fer est à la même température que le lac.

La variation d'entropie du fer est : $\Delta S_{\text{fer}} = m_{\text{fer}} \ln \frac{T_{\text{lac}}}{T_0} = -10,6 \cdot 10^6 \text{ J.K}^{-1}$.

L'entropie échangée est $S_{\text{éch}} = \frac{Q}{T_{\text{lac}}} = \frac{m_{\text{fer}} (T_{\text{lac}} - T_0)}{T_{\text{lac}}} = -19,9 \cdot 10^6 \text{ J.K}^{-1}$.

L'entropie créée est donc : $S_{\text{créée}} = \Delta S_{\text{fer}} - S_{\text{éch}} = 9,1 \cdot 10^6 \text{ J.K}^{-1} > 0$: la transformation est irréversible.

Cette création d'entropie est due la mise en contact de deux corps à températures différentes. La température de l'ensemble est inhomogène, il y a donc un transfert thermique du morceau de fer vers le lac, ce phénomène est irréversible donc crée de l'entropie.

A.2 1. Le système considéré est constitué des gaz dans les deux compartiments, de l'enceinte et du piston.

• Gaz (1) : état initial (P_1, V_1, T) , état final (P', V'_1, T') .

• Gaz (2) : état initial (P_2, V_2, T) , état final (P', V'_2, T') .

Le nombre de moles dans le compartiment de gauche est $n_2 = 2n_1$, puisque les volumes et températures sont égaux et que $P_2 = 2P_1$.

La paroi étant diatherme, la température finale est la même dans les deux compartiments. La cloison étant mobile, la pression finale est la même pour les deux gaz. On applique le premier principe pour déterminer l'état final :

$$\Delta U_1 + \Delta U_2 + \Delta U_{\text{enceinte}} + \Delta U_{\text{cloison}} + \Delta E_{\text{cloison}} \simeq \Delta U_1 + \Delta U_2 = W + Q$$

Le système est isolé thermiquement et l'enceinte est indéformable donc :

$$\Delta U_1 + \Delta U_2 = 0 \Leftrightarrow \frac{n_1 R}{\gamma - 1} (T' - T) + \frac{2n_1 R}{\gamma - 1} (T' - T) = 0 \Leftrightarrow T' = T$$

Pour déterminer les pressions et les volumes, on utilise l'équation d'état du gaz parfait et l'invariance du volume total, égal à $2V$. A l'état final, $P' V'_1 = n_1 R T$ et $P' V'_2 = n_2 R T$. En calculant le rapport des deux expressions : $\frac{V'_2}{V'_1} = \frac{n_2}{n_1} = 2$. Sachant que $V'_1 + V'_2 = V$, on en déduit $V'_1 = \frac{V}{3}$ et $V'_2 = \frac{2V}{3}$.

Pour la pression, les équations d'état donnent : $P' = \frac{P_1 + P_2}{2} = 1,5 \text{ bar}$.

2. Puisque l'évolution est adiabatique, l'entropie échangée est nulle et l'entropie créée est égale à la variation d'entropie totale des gaz :

$$S_{\text{créée}} = \Delta S_T = n_1 R \ln \frac{V'_1}{\frac{V}{2}} + n_2 R \ln \frac{V'_2}{\frac{V}{2}} = 2,83 \cdot 10^{-2} \text{ J.K}^{-1}$$

L'entropie créée est strictement positive ce qui prouve que la transformation est irréversible.

A.3 1. On écrit l'identité thermodynamique $dU = TdS - PdV$. Or pour un gaz parfait, on a $dU = C_V dT$.

Or la relation de Mayer donne $C_p - C_V = nR$ (on peut l'établir en écrivant l'identité thermodynamique pour l'enthalpie $dH = C_p dT = dU + d(PV) = C_V dT + nR dT$). Alors avec la définition du coefficient γ : $C_p = \gamma C_V$, on en déduit $C_V = \frac{nR}{\gamma - 1}$.

Par ailleurs, la relation $PV^k = \text{constante}$ donne en écrivant la dérivée logarithmique $\frac{dP}{P} + k \frac{dV}{V} = 0$ d'où $kPdV = -VdP$.

Or $d(PV) = nRdT = PdV + VdP = (1 - k)PdV$ donc $PdV = \frac{nR}{1 - k} dT$.

Finalement en reportant dans l'identité thermodynamique

$$dS = nR \left(\frac{1}{\gamma - 1} + \frac{1}{1 - k} \right) \frac{dT}{T}$$

2. Par intégration de la relation précédente, on obtient

$$\Delta S = nR \left(\frac{1}{\gamma - 1} + \frac{1}{1 - k} \right) \ln \frac{T_f}{T_i}$$

3. Comme la transformation est quasistatique, on peut identifier la pression extérieure à la pression du système. Le travail est donc $\delta W = -PdV = \frac{nR}{k - 1} dT$ en utilisant les relations établies à la première question. Par intégration, on a $W = \frac{nR}{k - 1} \Delta T$.

La variation d'énergie interne vaut $\Delta U = \frac{nR}{\gamma - 1} \Delta T$. Cela permet de calculer le transfert thermique

par l'expression du premier principe $Q = \Delta U - W$ donc $Q = nR \left(\frac{1}{\gamma - 1} + \frac{1}{1 - k} \right) \Delta T$.

4. i. Si $k = \gamma$, on obtient $Q = 0$ et l'évolution est adiabatique.

- ii. Si k tend vers l'infini, on a $W = 0$ donc la transformation est isochore.
- iii. Si $k = 1$, le travail tend vers l'infini, ce qui est impossible sauf si $\Delta T = 0$ et on a une transformation isotherme.
- iv. Si $k = 0$, le travail s'écrit $W = -nR\Delta T = -P_0\Delta V$ et la transformation est isobare.

A.4 1. La réponse est c.

Du fait de l'isolation thermique, on a $Q = 0$. D'autre part, $W = 0$ car les parois sont immobiles. Par conséquent, le premier principe s'écrit $\Delta U = 0$. En utilisant l'extensivité de l'énergie interne $n_1 c_v (T_f - T_1) + n_2 c_v (T_f - T_2) = 0$. On en déduit $T_f = \frac{n_1 T_1 + n_2 T_2}{n_1 + n_2}$.

2. La réponse est d.

D'après l'équation d'état $P_1 V_1 = n_1 R T_1$, $P_2 V_2 = n_2 R T_2$ et $P_f (V_1 + V_2) = (n_1 + n_2) R T_f$. On peut donc écrire $P_f = \frac{(n_1 + n_2) R T_f}{V_1 + V_2}$ soit en explicitant l'expression de T_f

$$P_f = \frac{(n_1 + n_2) R (n_1 T_1 + n_2 T_2)}{(n_1 + n_2) \left(\frac{n_1 R T_1}{P_1} + \frac{n_2 R T_2}{P_2} \right)} = \frac{P_1 P_2 (n_1 T_1 + n_2 T_2)}{n_1 T_1 P_2 + n_2 T_2 P_1}.$$

3. La réponse est b.

En utilisant les définitions de l'énoncé, on obtient $P_{f,1} (V_1 + V_2) = n_1 R T_f$ soit $P_{f,1} = \frac{n_1}{n_1 + n_2} P_f$. De même, on a $P_{f,2} = \frac{n_2}{n_1 + n_2} P_f$.

4. La réponse est a.

On utilise l'identité thermodynamique pour le gaz parfait soit $dS = n c_p \frac{dT}{T} - nR \frac{dP}{P}$. On l'intègre pour une mole et on obtient $\Delta S = c_p \ln \frac{T_f^2}{T_1 T_2} - R \ln \frac{P_f^2}{P_1 P_2} + 2R \ln 2$.

5. La réponse est d.

Avec les hypothèses proposées, on a $P_f = P_0$, $T_f = T_0$ et $\Delta S = 2R \ln 2$: il s'agit du paradoxe de Gibbs.

A.5 Le système étudié est la mole de cuivre.

1. Le système étant une phase condensée, son énergie interne ne dépend que de la température soit $dU = C dT$.

L'identité thermodynamique $dU = T dS - P dV$ s'écrit pour une phase condensée : $dU = T dS$, soit : $dS = \frac{C dT}{T}$. Cette équation s'intègre en $\Delta S = 3R \ln \frac{T_f}{T_i} = 2,20 \text{ J.K}^{-1}$.

L'entropie échangée s'écrit $\delta S_{\text{éch}} = \frac{C dT}{T_S}$ qui donne par intégration

$$S_{\text{éch}} = 3R \frac{T_f - T_i}{T_S} = 1,12 \text{ J.K}^{-1}$$

On en déduit l'entropie créée : $S_{\text{cr}} = \Delta S - S_{\text{éch}} = 1,08 \text{ J.K}^{-1}$.

Comme $S_{\text{cr}} > 0$, la transformation est irréversible. Cette création d'entropie est due à la mise en contact de deux corps à des températures différentes. La température de l'ensemble est inhomogène, il y a donc un transfert thermique du thermostat vers le cuivre, ce phénomène est irréversible donc crée de l'entropie.

2. La variation d'entropie se calcule comme dans le cas précédent donc on obtient $\Delta S = 3R \ln \frac{T_f}{T_i} = 2,20 \text{ J.K}^{-1}$.

L'entropie échangée est due aux transferts thermiques avec la résistance donc on peut l'écrire

$$S_{\text{éch}} = \frac{Q}{T_a} = \frac{-RI^2 \tau}{T_a} = 13,5 \text{ kJ.K}^{-1}.$$

L'entropie créée vaut donc $S_{\text{cr}} = \Delta S - S_{\text{éch}} = 13,5 \text{ kJ.K}^{-1}$.

On arrive à la même conclusion.

A.6 1. Il suffit de différencier l'expression proposée soit

$$dS = c_V \frac{dU}{U + \frac{a}{V}} + c_V \frac{-\frac{a dV}{V^2}}{U + \frac{a}{V}} + R \frac{dV}{V - b} = \frac{c_V V}{UV + a} dU + \left(\frac{R}{V - b} - \frac{a c_V}{aV + UV^2} \right) dV$$

2. L'identité thermodynamique s'écrit $dU = TdS - PdV$ soit $dS = \frac{dU}{T} + \frac{P}{T}dV$. Par identification des deux différentielles de l'entropie, on en déduit :

i. l'expression de l'énergie interne par $\frac{c_V V}{UV + a} = \frac{1}{T}$ soit

$$U = c_V T - \frac{a}{V}$$

ii. l'équation d'état de la substance par $\frac{R}{V - b} - \frac{a c_V}{aV + UV^2} = \frac{P}{T}$ soit

$$P = \frac{RT}{V - b} - \frac{a}{V^2}$$

On peut noter qu'il s'agit d'un gaz de Van der Waals.

3. L'énergie interne U et le volume V sont des grandeurs extensives tandis que la pression P et la température T sont intensives. On ajoute donc la quantité de matière pour assurer "l'homogénéité" du caractère extensif soit $U = n c_V T - \frac{n^2 a}{V}$ et $P = \frac{n r T}{V - n b} - \frac{n^2 a}{V^2}$.

On peut aussi simplement remarquer que dans l'expression pour une mole, il s'agit de l'énergie interne molaire $U_m = U/n$ et du volume molaire, $V_m = V/n$, ce qui redonne bien les expressions pour n moles.

4. La détente de Joule - Gay Lussac s'opère à énergie interne constante soit $\Delta U = 0$ (se reporter au chapitre sur les applications du premier principe pour la description de la détente et la démonstration de la conservation de l'énergie interne). Compte tenu de l'expression de l'énergie interne et du fait que c'est une fonction d'état, on en déduit $\Delta T = \frac{n a (V_2 - V_1)}{c_V V_1 V_2} = -1,3 \text{ K}$.

5. On déduit de l'expression donnée pour l'entropie la variation d'entropie :

$$\Delta S = c_V \ln \frac{T + \Delta T}{T} + R \ln \frac{V_2 - b}{V_1 - b} = 5,67 \text{ J.K}^{-1}.$$

6. On trouve que ΔT est indépendant de n , ce qui est normal du fait du caractère intensif de la température.

De même, il est normal que ΔS dépend de n puisque l'entropie est extensive.

7. On a tout pu déduire de la seule expression de l'entropie en fonction de l'énergie interne et du volume : $S(U, T)$ est une fonction caractéristique.

A.7 1. Puisque l'eau et la résistance formant le système (S) sont complètement décrits par leur température, les fonctions d'état U , H et S ne dépendent que de T , donc ces fonctions restent constantes pendant la transformation qui est isotherme : $\Delta U = 0$, $\Delta H = 0$ et $\Delta S = 0$.

Puisque l'ensemble est maintenu à la température T , c'est que (S) est en contact avec un thermostat à la même température, donc l'entropie échangée est :

$$S_{\text{éch}} = \frac{Q}{T}$$

On applique le premier principe pour déterminer Q (il s'agit d'un liquide et d'un solide, donc $\Delta U \simeq \Delta H$, on utilise indifféremment l'une ou l'autre de ces deux fonctions) :

$$\Delta H = W_{\text{élec}} + Q = 0 \Rightarrow Q = -W_{\text{élec}} = -R i^2 t = -2 \text{ kJ}$$

d'où $S_{\text{éch}} = \frac{Q}{T} = -6,82 \text{ J.K}^{-1}$. On en déduit l'entropie créée :

$$S_{\text{créée}} = \Delta S - S_{\text{éch}} = S_{\text{éch}} = 6,82 \text{ J.K}^{-1} > 0$$

L'origine de l'irréversibilité de la transformation est l'effet Joule. Il traduit la puissance cédée par le champ électromagnétique à la matière (Cf. cours de seconde année). Cette puissance est toujours positive dans le cas d'un conducteur, elle correspond à un phénomène irréversible. Au niveau microscopique, on peut le modéliser par un frottement exercé sur les électrons lors de leur déplacement.

2. L'évolution est maintenant adiabatique donc : $\Delta H = W_{\text{elec}}$, soit :

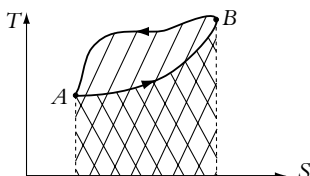
$$(m_e c_e + m_c c_c)(T_f - T_i) = RI^2 t$$

La température finale est $T_f = 297,7 \text{ K}$.

Puisque l'évolution est adiabatique, l'entropie créée est égale à la variation d'entropie :

$$S_{\text{créée}} = \Delta S_T = (m_e c_e + m_c c_c) \ln \frac{T_f}{T_i} = 6,77 \text{ J.K}^{-1} > 0$$

- A.8** 1. On considère la transformation réversible AB (chemin du bas) représentée sur le schéma ci-dessous. Elle est réversible donc le transfert thermique élémentaire s'écrit : $\delta Q = T dS$, et pour toute la transformation : $Q = \int_{AB} T dS$, ce qui représente l'aire sous la courbe. Pour la transformation BA (chemin du haut), le transfert thermique est négatif et est égal à l'opposé de l'aire sous la courbe.



Pour le cycle complet, le transfert thermique reçu par le système est donc égal, en valeur absolue, à l'aire du cycle. Dans l'exemple étudié, le cycle est décrit dans le sens trigonométrique et l'aire comptée négativement (correspondant à la transformation BA) est plus grande en valeur absolue que l'aire comptée positivement (correspondant à la transformation AB) : le transfert thermique reçu par le système au cours du cycle complet est négatif, le travail reçu est positif. Le cycle est récepteur. Pour qu'un cycle soit moteur, il faut qu'il soit décrit dans le sens horaire, comme en diagramme (P, V) .

2. On considère la transformation cyclique représentée par un des « carreaux » du maillage, décrit dans le sens horaire. L'aire du cycle est égale au travail fourni par le gaz. On représente maintenant les mêmes isothermes et les mêmes isentropiques dans le diagramme entropique (T, S) : les isothermes sont des droites horizontales et les isentropiques des droites verticales. Les aires de chaque rectangle ainsi dessinées sont égales et valent $T_0 S_0$. Or, comme on vient de la voir dans la question précédente, dans le diagramme entropique, l'aire d'un cycle représente le transfert thermique reçu par le gaz au cours du cycle. La transformation étudiée étant cyclique, $\Delta U = 0$ donc $|W| = |Q| = T_0 S_0$, ce qui répond à la question posée.

- B.1** 1. Le système étudié est la masse m d'eau. Pour un liquide, l'identité thermodynamique intégrée entre l'état initial et l'état final donne :

$$\Delta S = \int_{T_0}^{T_f} \frac{mc dT}{T} = mc \ln \frac{T_f}{T_0}$$

Pour calculer l'entropie échangée, il est nécessaire de calculer le transfert thermique reçu :

$$Q = \Delta H = mc(T_f - T_0)$$

et puisque le thermostat est à la température T_f , l'entropie échangée est :

$$S_{\text{éch}} = \frac{Q}{T_f} = mc \left(1 - \frac{T_0}{T_f} \right)$$

On en déduit l'entropie créée :

$$S_{\text{créée}} = \Delta S - S_{\text{éch}} = mc \left(\ln \frac{T_f}{T_0} - 1 + \frac{T_0}{T_f} \right) = 18 \text{ J.K}^{-1}$$

La cause de la création d'entropie est la différence de température entre la masse m et le thermostat.

2. Le système étudié est le même. La variation d'entropie de la masse d'eau est la même qu'à la question précédente puisque l'état initial et l'état final sont les mêmes. L'entropie échangée est la somme des entropies échangées avec chaque thermostat. On obtient :

$$S_{\text{créée}} = \Delta S - S_{\text{éch}} = mc \left(\ln \frac{T_f}{T_0} + \frac{T_0}{T_1} - 1 + \frac{T_1}{T_f} - 1 \right) = 9,22 \text{ J.K}^{-1}$$

L'entropie créée est donc environ deux fois moins importante que dans le cas d'un seul thermostat.

3. En procédant comme à la question précédente, on obtient :

$$\begin{aligned} S_{\text{créée}} &= \Delta S - S_{\text{éch}} = mc \ln \frac{T_f}{T_0} + mc \sum_{i=1}^N \left(\frac{T_{i-1}}{T_i} - 1 \right) \\ &= mc \ln \frac{T_f}{T_0} + mc \sum_{i=1}^N \left(\frac{1}{a} - 1 \right) \\ &= mc \left(\ln \frac{T_f}{T_0} + N \left(\frac{1}{a} - 1 \right) \right) \end{aligned}$$

Si N tend vers l'infini, a tend vers 0. Il faut exprimer a en fonction de N , T_0 et T_f . Sachant que

$a = \frac{T_i}{T_{i-1}}$, on en déduit que : $a^N = \frac{T_f}{T_0}$ donc que $a = \left(\frac{T_f}{T_0} \right)^{\frac{1}{N}}$, puis :

$$\begin{aligned} S_{\text{créée}} &= mc \left(\ln \frac{T_f}{T_0} + N \left(\frac{1}{a} - 1 \right) \right) \\ &= mc \left(\ln \frac{T_f}{T_0} + N \left(\left(\frac{T_f}{T_0} \right)^{-\frac{1}{N}} - 1 \right) \right) \\ &= mc \left(\ln \frac{T_f}{T_0} + N \left(\exp \left(-\frac{1}{N} \ln \frac{T_f}{T_0} \right) - 1 \right) \right) \\ &= mc \left(\ln \frac{T_f}{T_0} + N \left(1 - \frac{1}{N} \ln \frac{T_f}{T_0} - 1 \right) \right) \\ &= mc \left(\ln \frac{T_f}{T_0} - \ln \frac{T_f}{T_0} \right) \\ &= 0 \end{aligned}$$

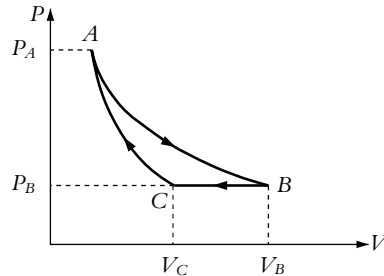
car, si N tend vers l'infini, $\frac{1}{N}$ tend vers 0 et on peut effectuer un développement limité au premier ordre en $\frac{1}{N}$ de l'exponentielle. Finalement, l'entropie créée tend vers 0. En multipliant le nombre de thermostats, on se rapproche d'une transformation réversible, l'entropie créée diminue jusqu'à s'annuler.

B.2 Le système étudié est le gaz qui décrit le cycle.

1. Pour calculer V_B , on applique l'équation du gaz parfait en B sachant que $T_B = T_A$, soit :

$$V_B = \frac{nRT_A}{P_B} = 24,9 \text{ L}$$

On a donc $V_B > V_C$. Le cycle tourne dans le sens horaire : c'est un cycle moteur.



2. On calcule la variation d'entropie entre A et B , en appliquant l'identité thermodynamique : $dU = T dS - P dV = 0$ pour une transformation isotherme d'un gaz parfait d'après la première loi

de Joule. On en déduit :

$$dS = \frac{P dV}{T} = nR \frac{dV}{V} \Rightarrow \Delta S_{AB} = nR \int_{V_A}^{V_B} \frac{dV}{V} = nR \ln \frac{V_B}{V_A} = nR \ln \frac{P_A}{P_B}$$

La transformation étant isotherme et le gaz parfait, $\Delta U_{AB} = 0$ donc $Q = -W$. L'évolution du gaz est de plus quasistatique, le travail reçu est donc $W = nRT_A \ln \frac{P_B}{P_A}$. Finalement :

$$Q = -nRT_A \ln \frac{P_B}{P_A}$$

On en déduit l'entropie échangée avec le thermostat à $T_T = T_A$:

$$S_{\text{éch}} = -nR \ln \frac{P_B}{P_A}$$

On remarque donc que $S_{\text{éch}} = \Delta S$: l'entropie créée est nulle ce qui prouve que la transformation est réversible.

3. Pour calculer la température en C, on applique encore l'équation du gaz parfait :

$$T_C = \frac{P_B V_C}{nR} = 246,6 \text{ K}$$

Il s'agit d'une transformation isobare donc :

$$W_{BC} = P_B(V_B - V_C) = 440 \text{ J}$$

et avec le premier principe :

$$Q_{BC} = \Delta H_{BC} = \frac{nR\gamma}{\gamma - 1} (T_C - T_B) = -1,55 \text{ kJ}$$

On peut aussi utiliser $Q_{BC} = \Delta U_{BC} - W_{BC}$, on obtient bien évidemment le même résultat.

On déduit du calcul précédent :

$$S_{\text{éch}} = \frac{Q}{T_T} = -5,17 \text{ J.K}^{-1}$$

Pour calculer l'entropie créée, il faut d'abord calculer la variation d'entropie du gaz entre B et C. Puisque c'est une évolution isobare, on utilise l'identité thermodynamique avec H : $dH = T dS + V dP = T dS$ car $dP = 0$, d'où :

$$dS = \frac{dH}{T} = \frac{nR\gamma}{\gamma - 1} \frac{dT}{T}$$

soit en intégrant :

$$\Delta S_{BC} = \frac{nR\gamma}{\gamma - 1} \ln \frac{T_C}{T_B} = -5,7 \text{ J.K}^{-1}$$

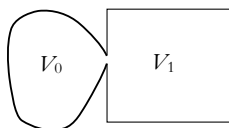
4. On peut alors calculer l'entropie créée :

$$S_{\text{créée}} = \Delta S - S_{\text{éch}} = -0,54 \text{ J.K}^{-1}$$

Ce résultat est en contradiction avec le second principe, puisque l'entropie créée ne peut être que positive : cette transformation est donc irréalisable.

L'entropie totale créée sur le cycle est celle créée entre B et C, puisqu'elle est nulle pour les deux autres transformations qui sont réversibles. Ce cycle est donc irréalisable pratiquement puisque l'entropie créée est négative. En revanche le cycle récepteur correspondant (ACBA) est possible. On verra effectivement dans le chapitre sur les machines thermiques qu'un cycle moteur monotherme (c'est-à-dire en contact avec un seul thermostat) ne peut exister.

- B.3** 1. On considère comme système l'air qui va entrer dans le récipient et le récipient lui-même dont on suppose la capacité thermique négligeable.



Le gaz entre dans le récipient jusqu'à ce que la pression soit égale à la pression extérieure P_0 en étant poussé par le gaz extérieur. Le gaz qui entre dans le récipient occupe, quand il est à l'extérieur, le volume V_0 . Il passe de la température T_0 (à l'extérieur) à la température T_1 (à l'intérieur). Le travail des forces de pression est le travail exercé par le gaz extérieur, et le volume balayé par le pourtour du gaz est V_0 . On en déduit que le travail reçu par le gaz est : $W = P_0 V_{\text{balayé}} = P_0 V_0$.
On applique le premier principe au système :

$$\Delta U = W + Q = W \Rightarrow \frac{nR}{\gamma - 1}(T_1 - T_0) = P_0 V_0$$

or $P_0 V_0 = nRT_0$, d'où $(T_1 - T_0) = (\gamma - 1)T_0$ et $T_1 = \gamma T_0$.

2. Comme l'évolution est adiabatique, l'entropie échangée est nulle et l'entropie créée est égale à la variation d'entropie du système. Pour calculer celle-ci, on retrouve rapidement l'expression de l'entropie d'un gaz parfait en variables (T, P) , en utilisant l'identité thermodynamique sous la forme $dH = T dS + V dP$ avec $dH = nC_{pm} dT = \frac{n\gamma R}{\gamma - 1} dT$ et $dP = 0$ puisque la pression initiale est égale à la pression finale. On obtient : $dS = \frac{n\gamma R}{\gamma - 1} \frac{dT}{T}$ soit :

$$S_{\text{créée}} = \Delta S = \frac{n\gamma R}{\gamma - 1} \ln \frac{T_1}{T_0} = \frac{nR\gamma}{\gamma - 1} \ln \gamma > 0$$

donc l'évolution est irréversible comme on pouvait s'y attendre.

L'origine de l'irréversibilité est la différence de densité entre l'intérieur et l'extérieur du récipient.

B.4 1. Dès que la fenêtre est ouverte, le gaz de l'enceinte A diffuse vers l'enceinte B . Si l'enceinte B est suffisamment grande (V_B supérieure à un volume critique appelé V_C), tout le gaz de A se retrouve dans B et dans l'état final, le piston est contre la paroi. Si $V_B < V_C$, il reste du gaz dans A et le piston s'immobilise avant la cloison. Dans ce cas, comme le piston est à l'équilibre dans l'état final, la pression dans le gaz est P_0 .

2. a) Le travail reçu par le gaz est dû au déplacement du piston, il est égal à $P_0 V_{\text{balayé}}$. Le volume balayé est :

$$V_{\text{balayé}} = V_A + V_B - V_1$$

donc le travail est :

$$W = P_0(V_A + V_B - V_1)$$

- b) Puisque $V_B < V_C$, la pression finale du gaz est P_0 . Le système (S) considéré est constitué du gaz, du cylindre et du piston.

L'état initial et l'état final du gaz sont :

$$\begin{array}{c|c} \text{Etat initial} & \begin{array}{l} P_0 \\ V_A \\ T_0 \end{array} \\ \hline \text{Etat final} & \begin{array}{l} P_1 = P_0 \\ V_1 \\ T_1 \end{array} \end{array}$$

L'application du premier principe à (S) donne :

$$\Delta U_{\text{gaz}} + \Delta U_{\text{enceinte}} + \Delta U_{\text{piston}} + \Delta E_{\text{cpiston}} = W + Q$$

On suppose que les capacités thermiques du piston et de l'enceinte sont négligeables, donc $\Delta U_{\text{piston}} \simeq 0$ et $\Delta U_{\text{enceinte}} \simeq 0$. Le piston est immobile dans les états extrêmes donc $\Delta E_{\text{cpiston}} = 0$. Enfin la transformation est adiabatique donc $Q = 0$. Finalement :

$$\Delta U = W \Rightarrow \frac{nR}{\gamma - 1}(T_1 - T_0) = P_0(V_A + V_B - V_1)$$

L'équation du gaz parfait donne $nRT_1 = P_0 V_1$. En combinant les deux expressions précédentes, on obtient pour T_1 :

$$T_1 = \frac{P_0}{nR} \left(\frac{\gamma - 1}{\gamma} V_B + V_A \right)$$

et pour V_1 :

$$V_1 = \frac{\gamma - 1}{\gamma} V_B + V_A$$

- c) Pour calculer la variation d'entropie du gaz, on imagine une transformation quasi-statique isobare allant du même état initial au même état final. Puisque S est une fonction d'état, sa variation est la même dans les deux transformations. On applique l'identité thermodynamique entre deux états d'équilibres voisins :

$$dS = \frac{nR\gamma}{\gamma-1} \frac{dT}{T} - \frac{VdP}{T} = \frac{nR\gamma}{\gamma-1} \frac{dT}{T} \Rightarrow \Delta S = \frac{nR\gamma}{\gamma-1} \ln \frac{T_1}{T_0}$$

► **Remarque :** On peut aussi utiliser (ou retrouver rapidement) l'expression de l'entropie d'un gaz parfait en variables (P, T) puisqu'on connaît la pression et la température dans l'état initial et dans l'état final.

On introduit les valeurs des températures trouvées précédemment :

$$\Delta S = \frac{nR\gamma}{\gamma-1} \ln \left(\frac{P_0}{nRT_0} \left(\frac{\gamma-1}{\gamma} V_B + V_A \right) \right) = \frac{nR\gamma}{\gamma-1} \ln \left(1 + \frac{\gamma-1}{\gamma} \frac{V_B}{V_A} \right)$$

Comme l'évolution est adiabatique, l'entropie échangée est nulle et donc l'entropie créée est égale à la variation d'entropie :

$$S_{\text{créée}} = \Delta S = \frac{nR\gamma}{\gamma-1} \ln \left(1 + \frac{\gamma-1}{\gamma} \frac{V_B}{V_A} \right) > 0$$

La transformation est bien irréversible. L'origine de l'irréversibilité est la différence de densité entre l'intérieur et l'extérieur du récipient A.

- d) A la limite, si $V_B = V_C$, le volume final V_1 est égal à V_B , soit :

$$\frac{\gamma-1}{\gamma} V_C + V_A = V_C \Rightarrow V_C = \gamma V_A$$

L'application numérique donne :

$$V_C = \gamma V_A = \gamma \frac{nRT_0}{P_0} = 0,035 \text{ m}^3 = 35 \text{ L}$$

3. Dans ce cas, on ne connaît pas la pression dans l'état final puisque le piston est retenu par la cloison, mais on connaît le volume final. L'état initial et l'état final du gaz sont :

$$\begin{array}{c|c} \text{Etat initial} & \text{Etat final} \\ \hline P_0 & P_1 \\ V_A & V_B \\ T_0 & T_1 \end{array}$$

Le volume balayé est V_A donc le travail est $W = P_0 V_A$. Le premier principe donne :

$$\Delta U_{\text{gaz}} = W \Rightarrow \frac{nR}{\gamma-1} (T_1 - T_0) = P_0 V_A$$

Or, $P_0 V_A = nRT_0$, donc :

$$T_1 = \gamma T_0 \quad \text{et} \quad P_1 = \frac{nRT_1}{V_B} = \gamma P_0 \frac{V_A}{V_B}$$

On remarque que $P_1 < P_0$, ce qui confirme la condition de validité de l'hypothèse $V_B > V_C$.

Pour calculer la variation d'entropie, on utilise (ou on retrouve rapidement) l'expression de l'entropie d'un gaz parfait en variables (T, V) (par exemple) :

$$\Delta S = \frac{nR}{\gamma-1} \ln \frac{T_1}{T_0} + nR \ln \frac{V_B}{V_A} = \frac{nR}{\gamma-1} \ln \frac{T_1 V_B^{\gamma-1}}{T_0 V_A^{\gamma-1}} = \frac{nR}{\gamma-1} \ln \left(\gamma \left(\frac{V_B}{V_A} \right)^{\gamma-1} \right)$$

L'évolution est adiabatique, donc l'entropie créée est égale à ΔS . On sait que $V_B > V_C$ donc $V_B > \gamma V_A$, et :

$$\frac{V_B}{V_A} > \gamma \Rightarrow \gamma \left(\frac{V_B}{V_A} \right)^{\gamma-1} > \gamma^{\gamma-1} > 1$$

L'argument du logarithme est plus grand que 1, $S_{\text{créée}} > 0$ et la transformation est irréversible, pour la même raison que précédemment.

4. a) Le système est le même que précédemment et puisque $V_B < V'_C$, le piston s'arrête avant la cloison donc la pression finale est P_0 . L'état initial et l'état final du gaz sont :

$$\text{Etat initial} \left| \begin{array}{c} P_0 \\ V_A \\ T_0 \end{array} \right. \quad \text{Etat final} \left| \begin{array}{c} P_0 \\ V_1 \\ T_0 \end{array} \right.$$

On en déduit le volume final : $V_1 = \frac{nRT_0}{P_0} = V_0$. Le piston n'a fait que transvaser le gaz sans changer son état.

- b) Le nouveau volume critique vérifie $V_1 = V'_C$, soit $V'_C = V_A$.

- c) Puisque les variables d'état du gaz sont inchangées, la variation d'entropie du gaz est nulle :

$$\Delta S = 0 \Rightarrow S_{\text{créée}} = -S_{\text{éch}}$$

On calcule l'entropie échangée avec le thermostat à la température T_0 :

$$S_{\text{éch}} = \frac{Q}{T_0}$$

On utilise le premier principe pour calculer Q . Sachant que $\Delta U = 0$ puisque T est constante :

$$Q = -W = -P_0 V_A = -nRT_0 \quad \text{donc : } S_{\text{éch}} = -nR \quad \text{et} \quad S_{\text{créée}} = nR > 0$$

La transformation est toujours irréversible.

B.5 1. L'équation d'état des gaz parfait est $PV = nRT$ où P est la pression en Pa, V le volume en m^3 , n la quantité de matière en mole, T la température en K et $R = 8,31 \text{ J.K}^{-1} \cdot \text{mol}^{-1}$ la constante des gaz parfaits.

2. Pour un gaz parfait, l'enthalpie et l'énergie interne ne dépendent que de la température et vérifient $H = U + PV = U + nRT$.

On en déduit $dH = nc_p dT = dU + d(PV) = nc_V dT + nR dT$ soit la relation de Mayer $c_p = c_V + R$.

3. Par définition de γ , on a $c_p = \gamma c_V$ donc en reportant dans la relation de Mayer, on en déduit $c_V = \frac{R}{\gamma - 1}$ et $c_p = \frac{\gamma R}{\gamma - 1}$.

4. On a $dU = nc_V dT$. Par l'identité thermodynamique $dU = TdS - PdV$ et la loi des gaz parfaits $PV = nRT$, on en déduit $dS = nc_V \frac{dT}{T} + nR \frac{dV}{V}$. Comme $H = U + PV$, on a aussi $dS = nc_p \frac{dT}{T} - nR \frac{dP}{P}$.

5. L'état final est défini par $T'_1 = T_1$ et $V'_1 = 2V_1$. De l'équation d'état, on en déduit $P'_1 = \frac{P_1}{2}$.

6. La transformation est quasistatique donc on peut identifier la pression à la pression extérieure, ce qui permet d'écrire $\delta W = -PdV = -nRT_1 \frac{dV}{V}$ en utilisant le fait que la transformation est isotherme. On a donc $W_1 = -nRT_1 \ln 2$.

7. La variation d'énergie interne est nulle puisque $\Delta U_1 = nc_V \Delta T = 0$ du fait du caractère isotherme de la transformation. On en déduit le transfert thermique $Q_1 = -W_1 = nRT_1 \ln 2$.

8. L'extérieur peut être considéré comme un thermostat donc l'entropie échangée vaut $S_{\text{éch}} = \frac{Q_1}{T_1} = nR \ln 2$.

9. D'après la relation de la question 4, la variation d'entropie vaut $\Delta S_1 = nR \ln 2$.

10. Des deux questions précédentes, on déduit l'expression de l'entropie créée $S_{\text{cr}} = \Delta S - S_{\text{éch}} = 0$. La transformation est donc réversible.

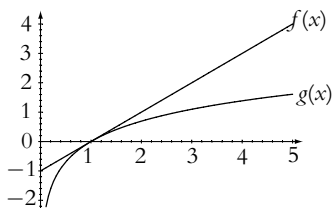
11. Les parois sont fixes donc il n'y a pas de travail $W = 0$. De plus, l'enceinte est calorifugée et il n'y a donc pas de transfert thermique $Q = 0$. On en déduit que la variation d'énergie interne est nulle $\Delta U_2 = 0$.

12. Comme les gaz sont parfaits, la nullité de la variation d'énergie interne entraîne l'absence de variation de température $\Delta T = 0$ soit $T_1 = T'_1$. L'état final vérifie d'autre part $V'_1 = 2V_1$ donc par la relation d'état, on en déduit $P'_1 = \frac{P_1}{2}$.
13. Pour un gaz non parfait, l'énergie interne U ne dépend pas uniquement de la température donc la variation de température ne serait pas nulle $\Delta T \neq 0$.
14. On a la même chose qu'au cas précédent puisque les états initial et final sont les mêmes.
15. L'entropie échangée est nulle $S'_{\text{éch}} = 0$ car le système est calorifugé.
16. On en déduit l'entropie créée $S'_{\text{cr}} = \Delta S_2 > 0$ donc la transformation est irréversible.
17. Les transformations correspondent aux mêmes états initial et final mais la transformation est réversible dans un cas et pas dans l'autre.

B.6 1. Le baromètre est une phase condensée, $\Delta S_B = C \ln \frac{T_0}{T_1} = -C \ln x$.

2. Pour le lac, $\Delta S_L = \frac{Q_L}{T_0}$ où Q_L est le transfert thermique reçu par le lac. Or l'ensemble {lac + baromètre} est isolé et n'échange que du transfert thermique donc $Q_L = -Q_B = -\Delta U_B = -C(T_0 - T_1)$. On en déduit $\Delta S_L = -\frac{C}{T_0}(T_0 - T_1) = C(x - 1)$.
3. Par extensivité de l'entropie, on en déduit la variation d'entropie totale

$$\Delta S = \Delta S_B + \Delta S_L = C(x - 1 - \ln x).$$
4. Le tracé de $f(x) = x - 1$ et $g(x) = \ln x$ donne :



La courbe représentative de g est sous celle de f pour toute valeur de x donc $\ln x \leq x - 1$ et $\Delta S \geq 0$.

5. Le premier principe s'écrit $\Delta U = \delta W + \delta Q_B + \delta Q_L = 0$ sur un cycle.
6. Le second principe dans le cas d'un échange thermique avec deux sources s'écrit $\frac{\delta Q_B}{T} + \frac{\delta Q_L}{T_0} \leq \Delta S$ avec $\Delta S = 0$ sur un cycle.
7. Dans le cas d'une transformation réversible l'inégalité précédente devient une égalité soit $\frac{\delta Q_B}{T} + \frac{\delta Q_L}{T_0} = 0$.
8. Le transfert thermique pour le baromètre s'écrit $\delta Q_B = dU_L = CdT$ donc $Q_B = C(T_0 - T_1)$. Alors par la relation de la question 7 et en utilisant l'hypothèse de réversibilité, on a : $\delta Q_L = CT_0 \frac{dT}{T}$. En intégrant, on obtient $Q_L = -CT_0 \ln \frac{T_0}{T_1}$.
9. Le rendement est défini comme le rapport entre ce qui est valorisé ici le travail et ce qui est dépensé ici le transfert thermique avec la source chaude c'est-à-dire le baromètre. On a donc $\eta = -\frac{W}{Q_B}$, le signe - permettant d'avoir un rendement positif.
10. D'après la relation issue du premier principe, le travail vaut

$$W = -Q_B - Q_L = C(T_1 - T_0) + CT_0 \ln \frac{T_0}{T_1} \quad \text{et le rendement s'écrit} \quad \eta = 1 - \frac{T_0}{T_0 - T_1} \ln \frac{T_0}{T_1}.$$

Chapitre 40

A.1 1. Si l'un de ces deux paramètres reste constant, l'autre aussi mais rien n'impose d'opérer à T ou à P constant. Par exemple, quand on introduit de l'eau liquide dans une enceinte calorifugée initialement vide, elle se vaporise et sa température diminue.

2. Pour une vaporisation isotherme par exemple, $\Delta U = mL_v - P_s \Delta V \neq 0$

3. $\Delta U = \Delta H - P_s \Delta V \neq \Delta H$.

4. Ce sont les paramètres intensifs qui sont fixés. En revanche, on ne sait rien sur la masse (ou le nombre de moles) du système dans chaque phase.

A.2 On considère comme système, le liquide à réfrigérer et le liquide réfrigérant. L'évolution est isobare (sous la pression atmosphérique) et rapide, donc on peut la supposer adiabatique. On exprime le premier principe à l'aide de la fonction enthalpie :

$$\Delta H_{\text{tot}} = Q \Rightarrow \Delta H_{\text{eau}} + \Delta H_{\text{refr}} = 0$$

On remplace les variations d'enthalpie par leurs expressions :

$$m_{\text{eau}} c (T_f - T_i) + m_{\text{refr}} L_v = 0 \Rightarrow T_f - T_i = -\frac{m_{\text{refr}} L_v}{m_{\text{eau}} c} = -8,7^\circ \text{C}$$

En pratique, le refroidissement est de l'ordre de 10°C en 5 s. L'ordre de grandeur trouvé est donc correct.

A.3 1. Soit m la masse de glace qui aura fondu et M_e la masse initiale d'eau. Dans l'état final pour lequel on a juste la quantité de glace nécessaire pour avoir de l'eau à $0,0^\circ \text{C}$ sans qu'il reste de glace, on a une masse d'eau $m + M_e$ et plus de glace. Le bilan énergétique de l'ensemble s'écrit : $\Delta H = 0$ soit $M_e c (\theta_f - \theta_i) + mL = 0$ donc la masse de glace nécessaire est

$$m = -\frac{M_e c (\theta_f - \theta_i)}{L} = 250 \text{ g}$$

Il s'agit de la masse minimale puisqu'au-delà de cette masse critique permettant d'avoir une température finale de $0,0^\circ \text{C}$, on a un mélange eau - glace à $0,0^\circ \text{C}$ qui est en équilibre et il n'y a donc plus d'évolution.

2. Pour l'eau, l'identité thermodynamique s'écrit : $dS_e = \frac{dU_e}{T} = M_e c \frac{dT}{T}$ (pas de variation de volume) et

$$\Delta S_e = M_e c \ln \frac{T_f}{T_i} = -297 \text{ J.K}^{-1}$$

3. Pour la glace se transformant en eau, on a $\Delta S_{\text{eg}} = \frac{mL}{T_0} = 308 \text{ J.K}^{-1}$.

4. Finalement la variation d'entropie pour l'ensemble est

$$\Delta S = \Delta S_e + \Delta S_{\text{ge}} = 11 \text{ J.K}^{-1}$$

Le système étant calorifugé, il n'y a pas d'entropie échangée et l'entropie créée est égale à la variation d'entropie soit $S_{\text{cr}} = \Delta S > 0$ et la transformation est irréversible.

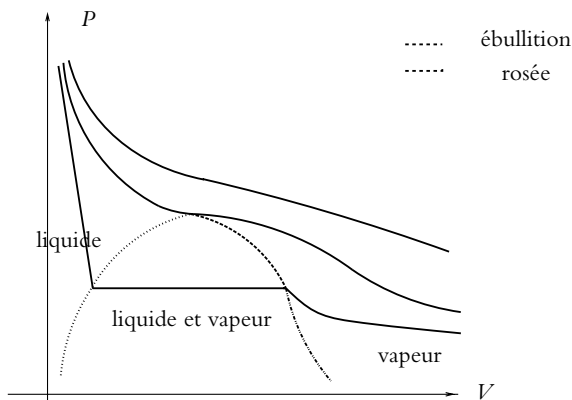
A.4 1. Le domaine (1) est le domaine de stabilité de la phase solide, (2) celui de la phase liquide et (3) celui de la phase gazeuse. Le point (B) est le point critique au-delà duquel on ne distingue plus les phases liquide et vapeur : on a une phase fluide.

Le point (E) est le point triple, en ce point coexistent les trois phases.

2. La pression de vapeur saturante est la pression d'équilibre entre la phase gazeuse et la phase liquide à une température donnée ; elle ne dépend que de la température.

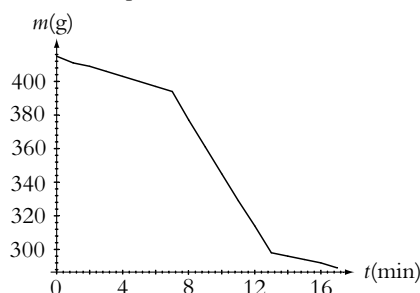
3. Le passage de la vapeur au liquide s'appelle liquéfaction.

4. Le diagramme demandé est le suivant :



5. Les vapeurs étant considérés comme des gaz parfaits, on peut écrire $P_s V = n_{\max} RT$. On en déduit $m_{\max} = \frac{M_{\text{eau}} P_s V}{RT}$.
6. Si $m < m_{\max}$, on a de la vapeur car cela correspond à $P < P_s$.
7. On n'a que du gaz si $P < P_s$ soit $V > V_{\text{lim}} = \frac{mRT}{M_{\text{eau}} P_s}$.
8. Le titre de vapeur est $x_{\text{vap}} = \frac{m_{\max}}{m} = \frac{M_{\text{eau}} P_s V}{mRT}$.

A.5 1. Le tracé de m en fonction du temps est le suivant :



On a donc trois segments de droite.

2. Dans la première phase, on a évaporation de l'azote ; dans la seconde, on ajoute le chauffage, ce qui accélère le phénomène et dans la troisième phase, on arrête le chauffage et on retrouve la première phase.
3. On estime la puissance de fuite thermique dans la phase 1 ou 3 soit $P_f = \left(\frac{dm}{dt} \right)_0 L_V$. On effectue ensuite le bilan lorsque le chauffage est branché, ce qui permet d'obtenir $P = P_f + UI$. On en déduit $L_V = \frac{UI}{\left(\frac{dm}{dt} \right) - \left(\frac{dm}{dt} \right)_0}$.
4. Le calcul des pentes donne $\left(\frac{dm}{dt} \right) = -16 \text{ g.min}^{-1}$ et $\left(\frac{dm}{dt} \right)_0 = -3 \text{ g.min}^{-1}$. On déduit de la relation de la question précédente $L_V = 224 \text{ J.g}^{-1}$.

- B.1** 1. a) On applique le premier principe à la goutte entre t et $t + dt$. Son volume ne varie pas donc elle ne reçoit que du transfert thermique. D'autre part, elle est liquide donc : $dU \simeq dH$, soit :

$$dU \simeq dH = \delta Q \Rightarrow \frac{4}{3}\pi R^3 \rho c dT = 4\pi R^2 h(T_a - T) dt \Rightarrow \frac{dT}{dt} + \frac{3h}{\rho R c} T = \frac{3h}{\rho R c} T_a$$

- b) La solution de l'équation précédente est : $T(t) = T_a + (T_1 - T_a) \exp\left(-\frac{t}{\tau}\right)$ avec $\tau = \frac{\rho R c}{3h}$. L'application numérique donne : $\tau = 4,3$ s et $t = \tau \ln \frac{T_1 - T_a}{T - T_a} = 3,9$ s.
2. a) On suppose d'après l'énoncé que l'eau liquide n'est pas complètement solidifiée, donc à l'état final la température du mélange est $T_f = 0^\circ\text{C}$. On imagine la suite de transformations suivante :

$$m \text{ liquide à } T_i = -5^\circ\text{C} \xrightarrow{(1)} m \text{ liquide à } T_f = 0^\circ\text{C}$$

puis :

$$m \text{ liquide à } T_f = 0^\circ\text{C} \xrightarrow{(2)} mx \text{ liquide} + m(1-x) \text{ glace à } T_f = 0^\circ\text{C}$$

On rappelle que la variation d'une fonction d'état est indépendante du chemin suivi. On applique le premier principe sur les deux transformations :

$$\Delta H = Q = 0 \Leftrightarrow mc(T_f - T_i) - m(1-x)L_f = 0 \Leftrightarrow x = 1 - \frac{c(T_f - T_i)}{L_f} = 0,94$$

- b) On applique le premier principe à la goutte, sachant qu'elle reste à $T = 0^\circ\text{C}$ pendant la solidification. On appelle dm la masse d'eau solidifiée pendant dt et dH la variation d'enthalpie du système :

$$dH = 4\pi R^2 h(T_a - T) dt \text{ et } dH = -dmL_f \Rightarrow dm = \frac{4\pi R^2 h(T - T_a)}{L_f} dt$$

soit en intégrant :

$$m(t) = \frac{4\pi R^2 h(T - T_a)}{L_f} t + m_0$$

La masse de liquide qu'il reste à solidifier est : $m(t) - m_0 = \frac{4}{3}\pi R^3 \rho x$, soit après simplification :

$$t = \frac{xRL_f\rho}{3h(T - T_a)} = 21,5 \text{ s}$$

- B.2** 1. Le système considéré est l'eau contenue dans le cylindre.

La première transformation est réversible. Puisque le déplacement du piston est lent, on peut supposer que la pression à l'intérieur du cylindre est maintenue en permanence égale à la pression atmosphérique. Dans ce cas, la transformation a lieu à température et pression constante. On peut alors écrire :

$$\Delta H = Q_1 = mL_v = 2,25 \text{ kJ} \text{ et } \Delta S = \frac{Q}{T} = \frac{\Delta H}{T} = 6 \text{ J.K}^{-1}$$

On calcule maintenant la variation d'énergie interne :

$$\Delta U = \Delta H - \Delta(PV) = \Delta H - P_s(V_{\text{final}} - V_{\text{initial}}) \simeq \Delta H - P_s V_{\text{final}} = 2,1 \text{ kJ}$$

La deuxième transformation n'est pas réversible, mais les états initial et final de l'eau sont les mêmes que pour la première transformation. Comme les variations des fonctions d'états ne dépendent que des états extrêmes, les variations de U , H et S sont les mêmes pour les deux transformations. Dans cette transformation, par contre, la pression n'est pas constante mais, puisque l'enceinte est supposée de volume constant, le travail des forces de pression est nul. En appliquant le premier principe, on trouve :

$$\Delta U = Q_2 \Rightarrow Q_2 = \Delta H - P_s V_{\text{final}} = ML_v - P_s V_{\text{final}}$$

On remarque sur ces deux exemples que le transfert thermique dépend du chemin suivi.

2. • Premier cas : $S_{\text{éch}} = \frac{Q_1}{T} = \frac{mL_v}{T}$ et $S_{\text{créée}} = \Delta S - S_{\text{éch}} = 0$, et on retrouve que la transformation est réversible.
- Deuxième cas : $S_{\text{éch}} = \frac{Q_2}{T} = \frac{mL_v - P_s V_{\text{final}}}{T}$ et $S_{\text{créée}} = \Delta S - S_{\text{éch}} = \frac{P_s V_{\text{final}}}{T}$, et on retrouve que la transformation est irréversible puisque l'entropie créée est positive.

B.3 1. Le système considéré est le fluide décrivant le cycle.

On se sert des isentropiques pour calculer les fractions massiques en vapeur x_B et x_C . On exprime les entropies des quatre points en fonction de l'entropie S_0 du fluide dans l'état 0 :

- $S_A = S_0 + mc_L \ln \frac{T_2}{T_1}$;
- $S_B = S_0 + mx_B \frac{L_v(T_1)}{T_1}$;
- $S_C = S_0 + mx_C \frac{L_v(T_1)}{T_1}$;
- $S_D = S_0 + mc_L \ln \frac{T_2}{T_1} + m \frac{L_v(T_2)}{T_2}$.

La transformation AB est isentropique donc $S_A = S_B$ et :

$$x_B = \frac{T_1}{L_v(T_1)} c_L \ln \frac{T_2}{T_1}$$

De même pour CD , $S_C = S_D$:

$$x_C = \frac{T_1}{L_v(T_1)} \left(c_L \ln \frac{T_2}{T_1} + \frac{L_v(T_2)}{T_2} \right)$$

2. Sur le cycle, la variation d'énergie interne est nulle, donc :

$$\Delta U = W + Q_{BC} + Q_{DA} = 0$$

puisque AB et CD sont adiabatiques. Les transformations DA et BC sont isobares donc :

$$Q_{BC} = \Delta H_{BC} = m(x_C - x_B)L_v(T_1) = -m \frac{T_1}{T_2} L_v(T_2) \quad \text{et} \quad Q_{DA} = \Delta H_{DA} = mL_v(T_2)$$

On en déduit :

$$W = -mL_v(T_2) \left(\frac{T_1}{T_2} - 1 \right) = 90,3 \text{ kJ.kg}^{-1}$$

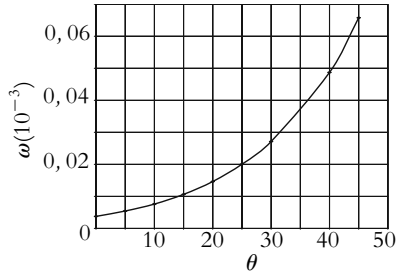
3. Le rendement trouvé est celui de Carnot : $\eta = \frac{1}{\frac{T_2}{T_1} - 1}$. Il s'agit effectivement d'un cycle de Carnot formé de deux isothermes et de deux adiabatiques.**B.4** 1. La relation des gaz parfaits s'écrit $P_a V = \frac{m_a}{M_a} RT$ donc par identification avec la relation donnée

$$R_a = \frac{R}{M_a} = 287 \text{ J.kg}^{-1}.\text{K}^{-1}.$$

2. De même, on obtient $R_v = \frac{R}{M_v} = 462 \text{ J.kg}^{-1}.\text{K}^{-1}$.3. On exprime $\omega = \frac{m_v}{m_a} = \frac{P_v R_a}{P_a R_v}$ en utilisant les relations précédentes. On obtient $\omega = A \frac{P_v}{P - P_v}$ avec $A = \frac{R_a}{R_v} = 0,62$.4. Par la définition de ϵ , on a $P_v = \epsilon P_{vs} = 1450 \text{ Pa}$ avec les valeurs du tableau. On en déduit $m_v = \frac{P_v V}{R_v T} = 10,9 \text{ g}$ d'après les questions précédentes. On en déduit également $m_a = \frac{(P - P_v) V}{R_a T} = 1208 \text{ g}$.5. De l'expression $\omega = A \frac{P_{vs}}{P - P_{vs}}$ avec $A = 0,62$, on établit le tableau de valeurs suivant :

$\theta(^{\circ}\text{C})$	0,0	5,0	10	15	20	25	30	40	45
$P_{vs} \text{ (Pa)}$	610	880	1227	1706	2337	3173	4247	7477	9715
$\omega(10^{-3})$	3,76	5,43	7,60	10,6	14,6	20,0	27,1	49,4	65,8

et la courbe :



On a du brouillard dans la zone au-dessus de la courbe.

6. On a $P_v = P_{vs}(10^\circ \text{C}) = 1227 \text{ Pa}$ donc $\epsilon = \frac{P_v}{P_{vs}(10^\circ \text{C})} = \frac{1227}{4247} = 0,29$.

7. Par sa définition, on a $h = \frac{H_a + H_v}{m_a + m_v} = \frac{m_a c_{pa} \theta + m_v c_{pv} \theta + m_v l}{m_a + m_v}$.

8. On obtient $H^* = h$ pour $m_a = 1$ et $m_v = \omega$ soit $H^* = \frac{c_{pa} \theta + \omega (c_{pv} \theta + l)}{1 + \omega}$.

9. Il n'y a pas de modifications des quantités de matière donc $\omega = \omega_q$.

La pression est constante donc $P_a + P_v = P$ est constant et $\frac{P_v}{P_a} = \frac{m_v R_v}{m_a R_a}$ également. On en déduit que P_a et P_v sont constantes.

Alors $Q = (m_a + m_v) \Delta h = (m_a c_{pa} + m_v c_{pv}) (\theta_q - \theta)$ et $\theta_q = \theta + \frac{(m_a + m_v) \Delta h}{m_a c_{pa} + m_v c_{pv}}$.

10. Comme P_v est constante, θ augmente ainsi que P_{vs} . On en déduit que ϵ diminue.

11. Comme à la question 6, $P_v = \epsilon P_{vs}(15^\circ \text{C}) = 0,85 * 1706 = 1450 \text{ Pa}$ et on en déduit $\omega = A \frac{P_v}{P - P_v} = 9,010^{-3}$.

La masse d'air humide traitée par seconde est $\frac{D}{3600}$ soit pour l'air sec $\frac{D}{3600(1 + \omega)}$. On en déduit la puissance $\mathcal{P} = \frac{D}{3600(1 + \omega)} (c_{pa} + \omega c_{pv}) (\theta_q - \theta) = 101 \text{ kW}$.

12. Il n'y a pas de perte de matière donc $\omega_3 = \frac{m_{v1} + m_{v2}}{m_{a1} + m_{a2}} = \frac{\omega_1 m_{a1} + \omega_2 m_{a2}}{m_{a1} + m_{a2}}$.

L'opération est adiabatique isobare ($\Delta H = 0$) : $m_{a1} H_1^* + m_{a2} H_2^* = (m_{a1} + m_{a2}) H_3^*$ donc $H_3^* = \frac{H_1^* m_{a1} + H_2^* m_{a2}}{m_{a1} + m_{a2}}$.

13. L'application numérique donne $m_{a1} = \frac{1}{1 + \omega_1} = 0,966$ et $m_{a2} = \frac{1}{1 + \omega_2} = 0,996$. En utilisant les relations des questions précédentes, on a $\omega_3 = 0,0192$, $H_1^* = 121 \text{ kJ.kg}^{-1}$, $H_2^* = 35,0 \text{ kJ.kg}^{-1}$ et $H_3^* = 77,3 \text{ kJ.kg}^{-1}$.

14. De l'expression de H^* , on déduit $\theta_3 = \frac{(1 + \omega_3) H_3^* - l \omega_3}{c_{pa} + c_{pv} \omega_3} = 29,5^\circ \text{C}$. Cette température est en dessous de la courbe de saturation.

15. En procédant de même, on obtient $\epsilon_4 = 1$ donc $P_{v4} = 7377 \text{ Pa}$ et $\omega_4 = 0,049$ avec $m_{a4} = 1,0 \text{ kg}$ ainsi que $\epsilon_5 = 0,2$ donc $P_{v5} = 0,2 * 880 = 176 \text{ Pa}$ et $\omega_5 = 1,10 \cdot 10^{-3}$ avec $m_{a5} = 1,0 \text{ kg}$.

En utilisant la relation de la question 12, on en tire $\omega_6 = 0,025$. Par conséquent, avec le même raisonnement qu'à la question 14, avec $H_4^* = 159 \text{ kJ.kg}^{-1}$ et $H_5^* = 7,8 \text{ kJ.kg}^{-1}$, on en déduit $H_6^* = 83,4 \text{ kJ.kg}^{-1}$ et $\theta_6 = 21,8^\circ \text{C}$.

Le point se trouve au-dessus de la courbe donc dans la zone de brouillard : la théorie n'est pas valide ici.

16. En appliquant la même technique, on a $\omega_6 = 0,025$ et $\theta_6 = 24,8^\circ\text{C}$. On lit ω sur la courbe de saturation soit $\omega = 0,019$. On en déduit $P_v = \frac{\omega P}{A + \omega} = 3012 \text{ Pa}$ et une masse $m_v = \frac{P_v V}{R_v T} = 21,9 \text{ g}$ d'eau liquide.

Chapitre 41

- A.1** 1. Le rendement du moteur étudié est $\rho = \frac{500}{1500} = \frac{1}{3}$. Le rendement du moteur de Carnot correspondant est $\rho_C = 1 - \frac{T_f}{T_c} = 1 - \frac{400}{650} = 0,385 > \rho$, ce qui normal.
2. Le second principe, appliqué au fluide sur un cycle, s'écrit :

$$\Delta S = 0 = S_{\text{éch}} + S_{\text{cr}} = \frac{Q_f}{T_f} + \frac{Q_c}{T_c} + S_{\text{cr}}$$

Or, $Q_c = 1500 \text{ J}$ et $Q_f = -W - Q_c$ d'après le premier principe, avec $W = -500 \text{ J}$, donc $Q_f = -1000 \text{ J}$. Finalement, $S_{\text{cr}} = 0,19 \text{ J.K}^{-1}$ par cycle. L'entropie créée est strictement positive : le moteur est irréversible.

3. Si la dépense est la même (c'est-à-dire si Q_c est la même), le premier principe appliqué à la machine réelle et à la machine de Carnot permet d'écrire : $W_{\text{réel}} - W_{\text{rév}} = Q_{f,\text{réel}} - Q_{f,\text{réel}}$. Or, d'après le second principe, $Q_{f,\text{réel}} = -\frac{T_f}{T_c} - T_f S_{\text{cr}}$ et $Q_{f,\text{rév}} = -\frac{T_f}{T_c}$, donc $W_{\text{réel}} = W_{\text{rév}} + T_f S_{\text{cr}}$. Les travaux sont négatifs et l'entropie créée positive donc $|W_{\text{réel}}| < |W_{\text{rév}}|$: pour une dépense identique, le moteur réel fournit moins de travail que le moteur réversible.

- A.2** 1. Puisque la masse M s'échauffe et que l'on utilise une pompe à chaleur, c'est que M joue le rôle de source chaude. Les notations sont celles de la figure ci-contre. Le premier principe appliqué à la machine sur un nombre entier de cycles donne : $\Delta U_{\text{mach}} = 0 = W + Q_1 + Q_2$. Le premier principe appliqué à la masse M donne : $\Delta U_M = Mc(T_f - T_i) = -Q_2$. Le second principe appliqué à la machine sur un nombre entier de cycles donne : $\Delta S_{\text{mach}} = 0$ L'entropie échangée par la machine est avec la source froide est

$$S_{\text{éch},1} = \frac{Q_1}{T_1}$$

et avec la source chaude :

$$S_{\text{éch},2} = \int \frac{\delta Q_2}{T} = \int \frac{-dU_M}{T} = -Mc \ln \frac{T_f}{T_i}$$

L'entropie créée est donc :

$$S_{\text{créée}} = \Delta S_{\text{mach}} - S_{\text{éch}} = 0 + Mc \ln \frac{T_f}{T_i} - \frac{Q_1}{T_1} = Mc \ln \frac{T_f}{T_i} - \frac{Mc(T_f - T_i) + W}{T_1}$$

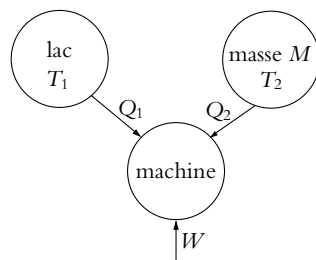
Puisque l'entropie créée est positive ou nulle, on tire de ces différentes relation :

$$W \geq Mc(T_f - T_i) - McT_1 \ln \frac{T_f}{T_i}$$

2. Application numérique : le travail minimal, correspondant au cas réversible, vaut $8,3 \cdot 10^6 \text{ J}$.
3. Si on avait utilisé une résistance chauffante directement, le premier principe appliqué à M donne :

$$\Delta U_M = W = Mc(T_f' - T_i) \Rightarrow T_f' - T_i = 2 \text{ K}$$

On aurait eu un échauffement de 2 K au lieu de 30 K d'où l'utilité de la pompe à chaleur.



A.3 1. La réponse est a.

On applique l'égalité de Clausius du fait du caractère réversible soit $\frac{\delta Q_c}{T_c} + \frac{\delta Q_f}{T_f} = 0$. On l'intègre

$$\frac{ML}{T_f} + \int_{T_c}^{T_{c,1}} 2MC \frac{dT_c}{T_c} = 0. \text{ Tout calcul fait, on obtient}$$

$$T_{c,1} = T_c \exp\left(-\frac{L}{2CT_f}\right) = 322 \text{ K}$$

2. La réponse est b.

On applique le premier principe sur un cycle $W + Q_c + Q_f = \Delta U = 0$ soit $W = -ML - 2MC(T_{c,1} - T_c) = -9,08 \text{ MJ}$.

3. La réponse est d.

On fait la même chose qu'en 1. soit $\int_{T_f}^{T_0} MC \frac{dT_f}{T_f} + \int_{T_{c,1}}^{T_0} 2MC \frac{dT_c}{T_c} = 0$ et finalement on trouve

$$T_0 = \sqrt[3]{T_f T_{c,1}^2} = 305 \text{ K}.$$

4. La réponse est d.

Comme à la question 2, on obtient $W' = -MC(T_0 - T_f) - 2MC(T_0 - T_{c,1}) - ML$ soit numériquement $W' = -10,2 \text{ MJ}$.

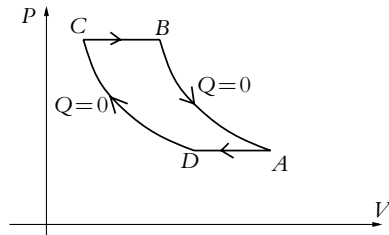
Le transfert thermique de la source chaude vaut $Q_c = 2MC(T_0 - T_c) = 56,8 \text{ MJ}$ et le rendement

$$\text{s'écrit } \eta = -\frac{W'}{Q_c} = 0,18.$$

5. La réponse est a.

On obtient par la démonstration du cours $\eta = 1 - \frac{T_f}{T_c} = 0,27$.

B.1 1.



2. Les relations de Laplace sont applicables si la transformation est adiabatique et quasistatique (ou réversible) et ne concernent que les gaz parfaits. L'expression de la relation de Laplace en fonction de la pression et de la température est $P^{1-\gamma} T^\gamma$ constante.

3. On applique la relation de Laplace à la transformation AB et on obtient $T_B = T_0 a^\beta = 448 \text{ K}$ puis à la transformation CD ce qui donne $T_D = T_1 a^{-\beta} = 188 \text{ K}$.

4. L'efficacité est définie par $e = -\frac{Q_{BC}}{W} = \frac{Q_{BC}}{Q_{BC} + Q_{DA}}$ en utilisant l'expression du premier principe $W + Q_{BC} + Q_{DA} = \Delta U = 0$ sur un cycle.

5. Les transformations sont isobares donc les transferts thermiques sont égaux à la variation d'enthalpie soit en tenant compte du caractère parfait des gaz $Q_{BC} = C_P(T_C - T_B)$ et $Q_{DA} = C_P(T_A - T_D)$.

On reporte dans l'expression de l'efficacité et on en déduit $e = \frac{1}{1 - a^{-\beta}} = 2,71$.

6. et 7. Le cycle de Carnot est composé de deux transformations adiabatiques et de deux transformations isothermes au cours desquelles ont lieu les transferts thermiques.

Le cycle est réversible donc on est dans le cas d'égalité de l'expression du second principe sur un cycle

$$\frac{Q_{BC}}{T_1} + \frac{Q_{DA}}{T_0} = 0. \text{ On en déduit } \frac{Q_{DA}}{Q_{BC}} = -\frac{T_0}{T_1} \text{ soit pour l'efficacité } e_C = \frac{T_1}{T_1 - T_0} = 19,9.$$

8. On a donc $e < e_C$, ceci s'explique par le caractère réversible du cycle de Carnot, ce qui n'est pas le cas *a priori* des autres cycles.

9. Le bilan d'entropie donne $\Delta S = S_{\text{éch}} + S_{\text{cr}} = 0$ sur le cycle. On a donc

$$S_{\text{cr}} = -S_{\text{éch}} = -\frac{Q_{BC}}{T_1} - \frac{Q_{DA}}{T_0} = \frac{nR}{\beta} \left(x + \frac{1}{x} - 2 \right)$$

10. L'étude de $f(x) = x + \frac{1}{x} - 2$ qui décroît de l'infini à 0 en $x = 1$ puis croît jusqu'à l'infini permet de conclure au signe positif de S_{cr} .

11. L'application numérique donne $S_{\text{cr}} = 4,84 \text{ J.K}^{-1}.\text{mol}^{-1}$.

12. Le régime est permanent donc ce qui est perdu est égal à ce qui est fourni. On en déduit que $\mathcal{P} = \frac{\dot{Q}_r}{e} = 7,4 \text{ kW}$.

B.2 1. En utilisant l'extensivité de l'entropie et le fait que le système global est calorifugé, on obtient :

$$\sigma = C \ln \frac{T_f}{T_1} + C \ln \frac{T_f}{T_2} \quad \text{avec} \quad T_f = \frac{T_1 + T_2}{2}.$$

2. On étudie une machine à trois sources. Les notations sont celles de la figure ci-contre. Le premier principe appliqué au fluide qui décrit des cycles donne :

$$\Delta U = W + Q_1 + Q_2 + Q_0$$

On applique le premier principe à Σ_1 , de températures initiale T_f et finale T_1 :

$$\Delta U_{\Sigma_1} = C(T_1 - T_f) = -Q_1$$

et pour Σ_2 , de températures initiale T_f et finale T_2 :

$$\Delta U_{\Sigma_2} = C(T_2 - T_f) = -Q_2$$

Le second principe appliqué au fluide qui décrit des cycles donne : $\Delta S = 0$. L'entropie échangée par ce système est :

$$S_{\text{éch}} = \int \frac{\delta Q_1}{T_{\Sigma_1}} + \int \frac{\delta Q_2}{T_{\Sigma_2}} + \frac{Q_0}{T_0} = -\int \frac{dU_{\Sigma_1}}{T_{\Sigma_1}} + \frac{Q_0}{T_0} - \int \frac{dU_{\Sigma_2}}{T_{\Sigma_2}} = -C \ln \frac{T_1}{T_f} - C \ln \frac{T_2}{T_f} + \frac{Q_0}{T_0}$$

L'entropie créée est donc :

$$S_{\text{créée}} = \Delta S - S_{\text{éch}} = C \ln \frac{T_1}{T_f} + C \ln \frac{T_2}{T_f} - \frac{Q_0}{T_0}$$

En remplaçant Q_0 dans l'expression précédente et sachant que l'entropie créée est positive ou nulle on en déduit :

$$W \geq T_0 C \ln \frac{T_f}{T_1} + T_0 C \ln \frac{T_f}{T_2} - C(T_1 - T_f) - C(T_2 - T_f)$$

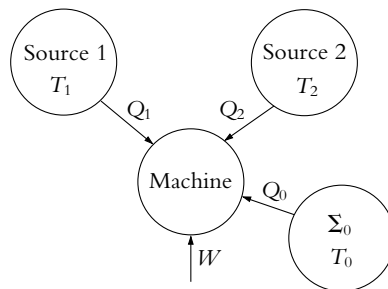
3. Pendant l'évolution décrite à la première question, l'entropie créée σ est strictement positive et égale à

$$\sigma = C \ln \frac{T_f}{T_1} + C \ln \frac{T_f}{T_2}$$

Si l'on utilisait une machine fonctionnant uniquement entre les deux sources Σ_1 et Σ_2 , l'entropie créée serait :

$$S_{\text{créée}} = C \ln \frac{T_1}{T_f} + C \ln \frac{T_2}{T_f} = -\sigma$$

L'évolution serait alors en désaccord avec le deuxième principe puisque l'entropie créée serait négative, le thermostat à température T_0 est donc nécessaire.



- B.3** 1. Dans tout l'exercice, on note T la température en kelvins et θ la température en degrés Celsius. On applique le premier et le deuxième principe au fluide qui décrit des cycles :

$$\begin{cases} W + Q_1 + Q_2 = 0 \\ \frac{Q_1}{T_1} + \frac{Q_2}{T_2} = 0 \end{cases}$$

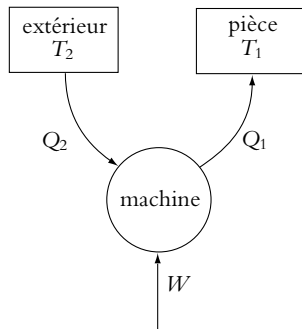
L'efficacité de la pompe à chaleur est :

$$e = \frac{-Q_1}{W} = \frac{T_1}{T_1 - T_2} = 13,3$$

Pour chauffer la pièce, il faut

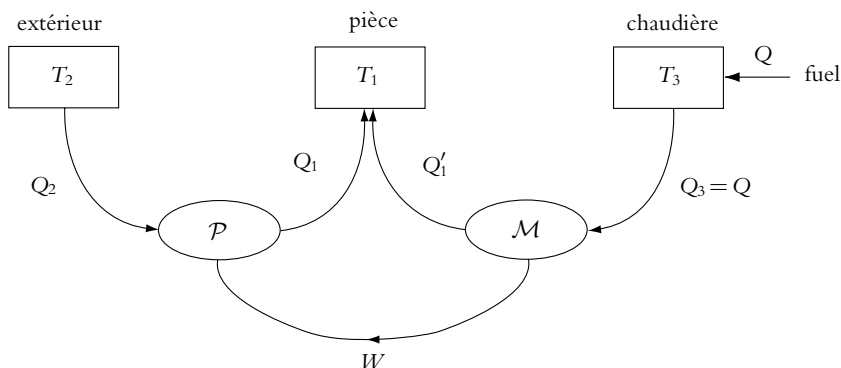
$$\dot{Q}_1 = 32 \text{ MJ.h}^{-1} = \frac{32}{3600} 10^6 \text{ W} = 8.9 \text{ kW}$$

$$\text{donc } \mathcal{P} = \frac{\dot{Q}_1}{e} = 668 \text{ W.}$$

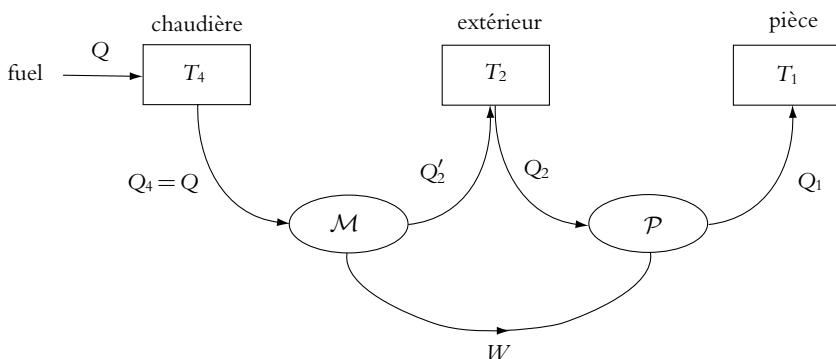


2. a) Les schémas de principe des deux dispositifs proposés sont les suivants ($\mathcal{P}$ désigne la pompe à chaleur et $\mathcal{M}$ le moteur) :

- Dispositif du premier conseiller :



- Dispositif du deuxième conseiller :



Dans chaque cas, on applique le premier principe et le second principe (sous forme de l'égalité de Clausius) au fluide qui décrit des cycles, pour la pompe à chaleur et pour le moteur. Attention, si on appelle W le travail reçu par le fluide au cours d'un cycle dans la pompe à chaleur, le travail que reçoit le fluide dans le moteur est $-W$.

• Dispositif du premier conseiller :

◦ Pompe à chaleur :

$$\begin{cases} W + Q_1 + Q_2 = 0 \\ \frac{Q_1}{T_1} + \frac{Q_2}{T_2} = 0 \end{cases}$$

◦ Moteur :

$$\begin{cases} -W + Q'_1 + Q = 0 \\ \frac{Q'_1}{T_1} + \frac{Q}{T_3} = 0 \end{cases}$$

La salle reçoit le transfert thermique :

$$Q_{\text{salle}} = -Q_1 - Q'_1 = \frac{1 - \frac{T_1}{T_3}}{1 - \frac{T_2}{T_3}} + \frac{T_1}{T_3} Q = \frac{1 - \frac{T_2}{T_3}}{1 - \frac{T_2}{T_1}} Q = \mathcal{P}_0 \Delta t_1$$

où $\mathcal{P}_0 = 32 \text{ MJ} \cdot \text{h}^{-1}$ et $Q = 32 \text{ MJ}$. On en déduit :

$$\Delta t_1 = \frac{1 - \frac{T_2}{T_3}}{1 - \frac{T_2}{T_1}} = 5 \text{ jours et 20 heures}$$

• Dispositif du deuxième conseiller :

◦ Pompe à chaleur :

$$\begin{cases} W + Q_1 + Q_2 = 0 \\ \frac{Q_1}{T_1} + \frac{Q_2}{T_2} = 0 \end{cases}$$

◦ Moteur :

$$\begin{cases} -W + Q + Q'_2 = 0 \\ \frac{Q}{T_4} + \frac{Q'_2}{T_2} = 0 \end{cases}$$

La salle reçoit le transfert thermique :

$$Q_{\text{salle}} = -Q_1 = \frac{1 - \frac{T_2}{T_4}}{1 - \frac{T_2}{T_1}} Q = \mathcal{P}_0 \Delta t_2$$

On en déduit :

$$\Delta t_2 = \frac{1 - \frac{T_2}{T_4}}{1 - \frac{T_2}{T_1}} = 6 \text{ jours et 13 heures}$$

- b) Le deuxième dispositif semble meilleur mais, si $T_3 = T_4$, les deux systèmes sont équivalents. Le premier semble plus astucieux car l'énergie thermique fournie à la source froide du moteur est utilisée pour chauffer la salle mais le rendement du moteur du deuxième dispositif est meilleur puisqu'il fonctionne entre la source la plus froide et la source la plus chaude.
- c) Si la température de la chaudière auxiliaire devient très grande, la durée de chauffage Δt tend vers $\frac{T_1}{T_1 - T_2}$, valeur que l'on obtient avec un moteur de rendement égal à 1. En fait, c'est l'efficacité de la pompe à chaleur qui limite la durée de chauffage. On ne peut pas avoir $\Delta t \rightarrow \infty$ conformément au second principe.

B.4 1. a) Lors d'une vaporisation totale d'une masse m d'un corps à température constante T , la variation d'enthalpie est :

$$\Delta H = mL_v = m(h_v - h_l)$$

L_v étant l'enthalpie massique de vaporisation. La variation d'entropie est :

$$\Delta S = \frac{mL_v}{T} = m(s_v - s_l)$$

- b) D'après ce qui précède, $s_v - s_l = \frac{h_v - h_l}{T}$, d'où $s_l(573) = 3,2 \text{ J.K}^{-1}$.
 c) L'équation des gaz parfaits écrite avec le volume massique donne :

$$Pv_v = \frac{RT}{M} \quad \text{soit} \quad v_v(293) = 58,8 \text{ m}^3.$$

2.

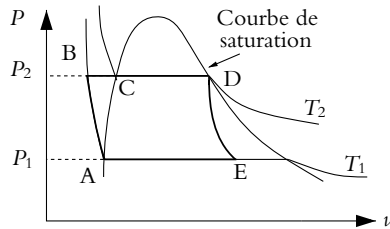


Figure S41.1 Diagramme de Clapeyron

3. a) La turbine est calorifugée donc la transformation est adiabatique. De plus l'énoncé suppose qu'elle est réversible. La transformation est donc adiabatique réversible : elle est isentropique.
 b) On utilise le fait que la transformation DE est isentropique, on calcule l'entropie du mélange liquide-vapeur dans ces deux états :
- en D, il n'y a que de la vapeur à 573 K, soit $S_D = ms_v(573)$;
 - en E, il y a une fraction x de vapeur à 293 K, soit

$$S_E = m(xs_v(293) + (1 - x)s_l(293))$$

$$\text{Soit avec } S_D = S_E, x = \frac{s_v(573) - s_l(293)}{s_v(293) - s_l(293)} = 0,68.$$

4. Dans l'alternateur le fluide est en écoulement horizontal. Dans ce cas, le travail W_{DE} reçu autre que celui des forces de pression est égal à la variation d'enthalpie. D'autre part, le travail reçu par l'alternateur est égal au travail cédé par le fluide soit :

$$W_a = -W_{DE} = H_D - H_E$$

Pour une masse m , $H_D = mh_v(573)$ et $H_E = m(xh_v(293) + (1 - x)h_l(293))$, et par unité de masse de fluide écoulée :

$$W_a = 1,140.10^6 \text{ J.kg}^{-1}$$

5. a) Entre B et D l'évolution est isobare, donc pour $m = 1 \text{ kg}$:

$$\Delta H_{BD} = Q_{BD} \Rightarrow Q_{BD} = c(T_2 - T_1) + h_v(573) - h_l(573) = 2,78.10^6 \text{ J.kg}^{-1}$$

- b) L'efficacité est : $e = \frac{1,14}{2,86} = 0,4$. L'énoncé utilise le terme efficacité, mais il s'agit plutôt d'un rendement.

Dans cette définition, on omet de tenir compte du travail nécessaire pour faire fonctionner la pompe, le rendement est donc légèrement plus faible mais ce travail est négligeable devant W_a .

- c) Le rendement maximal qu'on aurait pu avoir est celui de Carnot pour un moteur réversible :
- $$e_c = 1 - \frac{T_1}{T_2} = 0,49.$$

Le rendement obtenu est inférieur, donc certaines transformations ne sont pas réversibles, ce sont AB et BC.

- B.5** 1. En partant du point 1, le diazote à l'état gazeux subit une compression isotherme jusqu'à la pression P_2 qui est égale à P_3 puisque la transformation suivante est isobare. Le volume massique v_2 est inférieur à v_1 puisqu'il s'agit d'un gaz. Le point 2 est à l'intersection de l'isotherme T_2 et de l'isobare P_3 .
 Le point 3 est à l'horizontale du point 2 (transformation isobare) à l'intersection de l'isobare P_3 et de l'isotherme T_3 .

Le point 4 correspondant à un mélange liquide-vapeur à T_4 est sur l'isotherme T_4 entre les points A et B .

Le point 6 est le point A et le point 5 est le point B .

On obtient donc le diagramme ci-dessous :

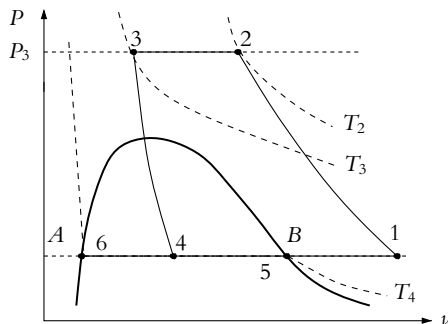


Figure S41.2

- Par définition, l'enthalpie massique de vaporisation à 78 K est $L_v = h_B - h_A$. Mais l'énoncé ne donne que s_A et s_B . On utilise alors la relation : $L_v = T_4(s_B - s_A)$ et on trouve $L_v = 196,56 \text{ kJ.kg}^{-1}$.
- La transformation 3-4 est une détente adiabatique réversible donc $s_4 = s_3$. Au point 4, on a un mélange liquide-vapeur de titre massique en vapeur noté x_4 , d'où :

$$s_3 = s_4 = x_4 s_B + (1 - x_4) s_A = x_4 (s_B - s_A) + s_A$$

On en déduit :

$$x_4 = \frac{s_3 - s_A}{s_B - s_A} = 0,52$$

- Comme pour l'entropie, on peut écrire :

$$h_4 = x_4 h_B + (1 - x_4) h_A = x_4 (h_B - h_A) + h_A = x_4 (h_5 - h_6) + h_5$$

- Comme l'énoncé l'indique, on considère le système (S). On s'inspire de la démonstration faite dans le cours pour un fluide en écoulement. Les pressions sont toutes les mêmes puisque l'évolution est isobare. Les masses dm_2 et dm_3 sont égales (notées dm). Il en est de même pour dm_5 et dm_1 (notées $dm' = x_4 dm$). Pour une masse de un kilogramme circulant dans l'échangeur entre l'état 2 et l'état 3, il n'y a plus que x_4 kilogramme circulant dans l'échangeur entre 5 et 1.

On applique le premier principe à (S) :

$$dU = dU_{\Sigma} + dU_{\Sigma'} + U_{dm_3} - U_{dm_2} + U_{dm_1} - U_{dm_5} = \delta Q + \delta W$$

Comme le régime est stationnaire $dU_{\Sigma} = dU_{\Sigma'} = 0$. La transformation est adiabatique donc $\delta Q = 0$. Le travail est celui des forces de pression :

$$\delta W = P_3 dm v_2 - P_3 dm v_3 + P_5 dm' v_1 - P_5 dm' v_5$$

En reportant cette expression dans celle du premier principe, on trouve :

$$dm(h_3 - h_2) + dm'(h_1 - h_5) = 0 \Rightarrow h_3 - h_2 = x_4(h_5 - h_1)$$

On en déduit :

$$h_1 = h_5 + \frac{h_2 - h_3}{x_4}$$

l'énoncé ne donne pas h_5 mais h_6 . On calcule h_5 grâce à la relation :

$$h_5 = h_6 + L_v = 231 \text{ kJ.kg}^{-1}$$

et on trouve :

$$h_1 = 629 \text{ kJ.kg}^{-1}$$

6. La compression 1-2 est isotherme réversible donc on peut écrire :

$$\Delta S = \int_1^2 \frac{\delta Q}{T_2} = \frac{Q}{T_2}$$

soit, pour une masse égale à 1 kg : $Q = T_2(s_2 - s_1) = -498,8 \text{ kJ.kg}^{-1}$

7. En appliquant le premier principe pour un fluide en écoulement, on obtient le travail reçu W_c par un kilogramme de fluide de la part du compresseur :

$$W_c = \Delta h - Q = h_2 - h_1 - Q = 290 \text{ kJ.kg}^{-1}$$

Une machine de 100 kW, consomme un travail de $3,6 \cdot 10^8 \text{ J}$ par heure et donc comprime par heure une masse m de diazote :

$$m = \frac{3,6 \cdot 10^8}{290 \cdot 10^3} = 721 \text{ kg}$$

Pour avoir la masse d'azote liquide produite, on multiplie la masse m par $(1 - x_4)$ ce qui donne 346 kg.h^{-1} .

8. En appliquant le premier principe pour un fluide en écoulement, le travail reçu par un kilogramme de fluide lors de la détente adiabatique de la part de la turbine est :

$$W_T = h_4 - h_3 = h_6 + x_4 L_v - h_3 = -60 \text{ kJ.kg}^{-1}$$

B.6 1. On a un cycle moteur s'il fournit du travail donc avec la convention usuelle de compter positivement ce qui est reçu, on doit avoir $W < 0$ pour un moteur. Dans ces conditions, l'interprétation graphique du travail impose d'avoir un sens horaire pour parcourir le cycle.

2. La transformation BC est isochore donc le travail est nul $W_{BC} = 0$ et le transfert thermique est $Q_{BC} = \Delta U_{BC} = m C_V (T_C - T_B)$.

La transformation CD est isobare donc le calcul du travail donne $W_{CD} = \Delta (PV)_{CD}$ et on en déduit celui du transfert thermique $Q_{CD} = \Delta H_{CD} = m C_P (T_D - T_C)$.

La transformation EA est isochore donc le travail est nul $W_{EA} = 0$ et le calcul du transfert thermique donne $Q_{EA} = \Delta U_{EA} = m C_V (T_A - T_E)$.

3. Les transferts thermiques sont tels que $Q_{BC} > 0$ et $Q_{CD} > 0$ donc ils correspondent à une dépense.

Le travail W est utile. Par conséquent, le rendement est défini par $r_{th} = -\frac{W}{Q_{BC} + Q_{CD}}$.

4. L'écriture du premier principe sur le cycle donne $\Delta U = 0$ soit en explicitant la variation d'énergie interne $\Delta U = W + Q_{AB} + Q_{BC} + Q_{CD} + Q_{DE} + Q_{EA}$. Par ailleurs, on a $Q_{AB} = Q_{DE} = 0$ car ces transformations sont adiabatiques. On en déduit $r_{th} = 1 + \frac{Q_{EA}}{Q_{BC} + Q_{CD}} = 1 + \frac{T_A - T_E}{\gamma(T_D - T_C) + T_C - T_B}$.

5. La transformation AB est adiabatique et quasistatique. De plus, elle concerne un gaz parfait donc on peut appliquer les relations de Laplace par exemple $TV^{\gamma-1}$ constant. On en déduit $T_B = a^{\gamma-1} T_A$.

La transformation BC est isochore donc en utilisant la loi des gaz parfaits et la relation précédente, on a $T_C = \frac{P_C}{P_B} T_B = d a^{\gamma-1} T_A$.

La transformation CD est isobare donc en utilisant la loi des gaz parfaits et la relation précédente, on obtient $T_D = \frac{V_D}{V_C} T_C = c d a^{\gamma-1} T_A$.

La transformation DE est adiabatique et quasistatique et elle s'applique à un gaz parfait donc les relations de Laplace sont valables comme $TV^{\gamma-1}$ constant soit $T_E = c^{\gamma} d T_A$.

6. En reportant les expressions de la question précédente, on obtient

$$r_{th} = 1 + \frac{1 - c^{\gamma} d}{(d-1)\gamma + (c-1)d} a^{1-\gamma}$$

7. La relation des gaz parfaits avec $n = \frac{m}{M_{air}}$ permet d'écrire $m = \frac{M_{air} P_A V_A}{R T_A}$ ainsi que

$$B = V_A - V_B = V_A \frac{a-1}{a}. \text{ On en déduit } m = \frac{a}{a-1} \frac{M_{air} P_A B}{R T_A}.$$

8. Par définition, $Q_{BC} = \eta \mu_V = \frac{k_V}{100} \eta \alpha m$.

9. On égale cette expression avec $Q_{BC} = m C_V (T_C - T_B)$, ce qui permet d'obtenir $T_C = a^{\gamma-1} T_A + \frac{k_V \eta \alpha}{100 C_V}$.

La relation des gaz parfaits donne $P_C = \frac{V_A}{V_C} \frac{T_C}{T_A} P_A$ soit en reportant les résultats précédents $P_C = a \left(a^{\gamma-1} + \frac{k_V \eta \alpha}{100 C_V T_A} \right) P_A$.

10. Maintenant on a $Q_{CD} = \eta (\mu - \mu_V) = \left(1 - \frac{k_V}{100} \right) \eta \alpha m$.

11. En égalant cette expression avec $Q_{CD} = m C_P (T_D - T_C)$, on obtient

$$T_D = a^{\gamma-1} T_A + \frac{\eta \alpha}{\gamma C_V} \left(\frac{k_V}{100} (\gamma - 1) + 1 \right)$$

12. On veut $P_C < P_{\max}$. A partir de l'expression de P_C , on en déduit qu'il faut vérifier $k_V < \frac{100 C_P T_A}{\eta \alpha \gamma} \left(\frac{P_{\max}}{P_A} - a^{\gamma} \right)$ soit $k_V < 21,4$.

13. Les applications numériques donnent :

- $V_A = \frac{a}{a-1} B = 3,21 \text{ L}$,
- $V_B = \frac{V_A}{a} = 0,214 \text{ L}$,
- $T_B = a^{\gamma-1} T_A = 972 \text{ K}$,
- $T_C = a^{\gamma-1} T_A + \frac{k_V \eta \alpha}{100 C_V} = 1464 \text{ K}$,
- $T_D = a^{\gamma-1} T_A + \frac{\eta \alpha}{\gamma C_V} \left(\frac{k_V}{100} (\gamma - 1) + 1 \right) = 2765 \text{ K}$,
- $c = \frac{T_D}{T_C} = 1,89$,
- $d = \frac{T_C}{T_B} = 1,51$,
- $T_E = c^{\gamma} d T_A = 1236 \text{ K}$,
- $b = \left(\frac{T_D}{T_E} \right)^{\frac{1}{\gamma-1}} = 7,88$.

14. On en déduit le rendement théorique $r_{th} = 1 + \frac{1 - c^{\gamma} d}{(d-1)\gamma + (c-1)d} a^{1-\gamma} = 0,55$.

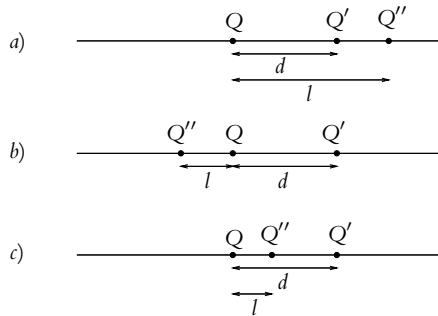
15. De même, on a $m = \frac{a}{a-1} \frac{M_{\text{air}} P_A B}{R T_A} = 3,31 \text{ g}$, ce qui permet de calculer $Q_{BC} = \frac{k_V}{100} \eta \alpha m = 1136 \text{ J}$ et $Q_{CD} = \left(1 - \frac{k_V}{100} \right) \eta \alpha m$. On en déduit $W = r_{th} (Q_{BC} + Q_{CD}) = 3073 \text{ J}$ par cycle et par cylindre. P est le produit du nombre de cylindres 6, du nombre de cycles par tour 0,5, du nombre de tours par seconde $\frac{2300}{60}$ et du travail soit $P = 353 \text{ kW}$.

16. Le cycle est réel et non réversible donc la puissance réelle est inférieure.

Partie VII – Électromagnétisme

Chapitre 43

- A.1** 1. La norme du champ électrostatique créé par Q' en M est : $E(M) = \frac{|Q'|}{4\pi\epsilon_0 d^2}$ et celle du champ créé par Q en M' : $E(M') = \frac{|Q|}{4\pi\epsilon_0 d^2}$. On a donc $E(M) > E(M')$ si $|Q'| > |Q|$, $E(M) < E(M')$ si $|Q'| < |Q|$ et $E(M) = E(M')$ si $|Q'| = |Q|$.
2. La norme de la force électrostatique subie par Q est : $F(M) = \frac{|QQ'|}{4\pi\epsilon_0 d^2}$ et celle subie par Q' : $F(M') = \frac{|QQ'|}{4\pi\epsilon_0 d^2}$. Les forces électrostatiques sont donc égales en module quelles que soient les valeurs respectives de Q et de Q' . De plus, $\vec{F}(M) = -\vec{F}(M')$ conformément au principe des actions réciproques.
3. On est dans l'un des trois cas suivants :



On applique le principe de superposition pour calculer les champs électrostatiques en M et en M' . Du fait des positions relatives des charges, le champ est toujours dirigé le long de l'axe MM' .

- cas a) : en projection sur l'axe MM' , le champ électrostatique en M vaut $E(M) = \frac{1}{4\pi\epsilon_0} \left(\frac{Q'}{d^2} + \frac{Q''}{l^2} \right)$ et en M' $E(M') = \frac{1}{4\pi\epsilon_0} \left(\frac{Q}{d^2} + \frac{Q''}{(l-d)^2} \right)$, on veut $E(M) = E(M')$ donc :

$$\frac{Q'}{d^2} + \frac{Q''}{l^2} = \frac{Q}{d^2} + \frac{Q''}{(l-d)^2}$$

- cas b) : en projection sur l'axe MM' , le champ électrostatique en M vaut $E(M) = \frac{1}{4\pi\epsilon_0} \left(\frac{Q'}{d^2} + \frac{Q''}{l^2} \right)$ et en M' $E(M') = \frac{1}{4\pi\epsilon_0} \left(\frac{Q}{d^2} + \frac{Q''}{(l+d)^2} \right)$, on veut $E(M) = E(M')$ donc :

$$\frac{Q'}{d^2} + \frac{Q''}{l^2} = \frac{Q}{d^2} + \frac{Q''}{(l+d)^2}$$

- cas c) : en projection sur l'axe MM' , le champ électrostatique en M vaut $E(M) = \frac{1}{4\pi\epsilon_0} \left(\frac{Q'}{d^2} + \frac{Q''}{l^2} \right)$ et en M' $E(M') = \frac{1}{4\pi\epsilon_0} \left(\frac{Q}{d^2} + \frac{Q''}{(d-l)^2} \right)$, on veut $E(M) = E(M')$ donc :

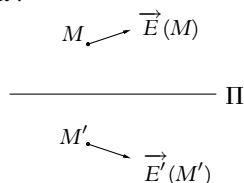
$$\frac{Q'}{d^2} + \frac{Q''}{l^2} = \frac{Q}{d^2} + \frac{Q''}{(d-l)^2}$$

On note qu'on trouve la même expression dans le cas a) et dans le cas c).

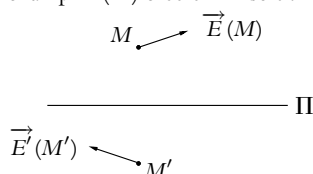
4. On se place dans le cas b). On déduit de l'égalité de la question précédente l'expression suivante pour

$$Q'' : Q'' = \frac{(Q - Q') l^2 (d - l)^2}{d^3 (d - 3l)}.$$

A.2 1. Le champ $\vec{E}'(M')$ en M' , symétrique de M par rapport à Π , est égal au symétrique par rapport à Π du champ $\vec{E}(M)$ créé en M soit :

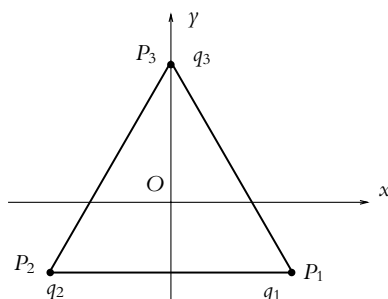


2. Le champ $\vec{E}'(M')$ en M' , symétrique de M par rapport à Π , correspond au champ créé par la distribution de charges opposées puisque Π est un plan d'antisymétrie. Il est donc égal à l'opposé du symétrique par rapport à Π du champ $\vec{E}(M)$ créé en M soit :



C'est aussi le symétrique de $\vec{E}(M)$ par rapport à la normale au plan Π .

B.1



On applique le principe de superposition :

$$\vec{E}(O) = \frac{1}{4\pi\epsilon_0} \left(\frac{q_1}{P_1O^3} \vec{P_1O} + \frac{q_2}{P_2O^3} \vec{P_2O} + \frac{q_3}{P_3O^3} \vec{P_3O} \right)$$

Or $P_1O = P_2O = P_3O = OP$ car O est le centre de gravité du triangle équilatéral $P_1P_2P_3$ donc :

$$\vec{E}(O) = \frac{1}{4\pi\epsilon_0 OP^3} (q_1 \vec{P_1O} + q_2 \vec{P_2O} + q_3 \vec{P_3O}).$$

1. Si $q_1 = q_2 = q_3 = q$ alors $\vec{E}(O) = \vec{0}$.

2. Si $q_1 = q_2 = q$ et $q_3 = -q$ alors $\vec{E}(O) = \frac{q}{4\pi\epsilon_0 OP^3} (\vec{P_1O} + \vec{P_2O} - \vec{P_3O})$. Or $\vec{P_1O} + \vec{P_2O} = \vec{OP_3}$ donc $\vec{P_1O} + \vec{P_2O} - \vec{P_3O} = 2\vec{P_3O}$. On en déduit : $\vec{E}(O) = \frac{q}{2\pi\epsilon_0 OP^3} \vec{u}_y$. On peut également obtenir ce résultat en considérant la superposition du champ créé par trois charges égales $q_1 = q_2 = q_3 = q$ (qui est nul d'après la première question) et du champ créé par une charge $q_3 = -q$. On obtient bien le champ créé par une charge $-q$ placée en P_3 .

3. Si $q_1 = q_2 = q$ et $q_3 = 2q$ alors $\vec{E}(O) = \frac{q}{4\pi\epsilon_0 OP^3} (\vec{P_1O} + \vec{P_2O} + 2\vec{P_3O})$. Or $\vec{P_1O} + \vec{P_2O} + 2\vec{P_3O} = \vec{P_3O}$ donc $\vec{E}(O) = -\frac{q}{4\pi\epsilon_0 OP^3} \vec{u}_y$. On peut également obtenir ce résultat en considérant la superposition du champ créé par trois charges égales $q_1 = q_2 = q_3 = q$ (qui est nul d'après la première question) et du champ créé par une charge $q_3 = q$. On obtient bien le champ créé par une charge q placée en P_3 .
4. Si $q_1 = q$, $q_2 = 2q$ et $q_3 = -q$ alors par projection sur les axes Ox et Oy , on en déduit les composantes du champ $\vec{E}(O)$:

$$\begin{cases} E_x = \frac{q}{4\pi\epsilon_0 OP^3} \left(\cos \frac{5\pi}{6} + 2 \cos \frac{\pi}{6} - \cos \frac{-\pi}{2} \right) = \frac{q\sqrt{3}}{8\pi\epsilon_0 OP^3} \\ E_y = \frac{q}{4\pi\epsilon_0 OP^3} \left(\sin \frac{5\pi}{6} + 2 \sin \frac{\pi}{6} - \sin \frac{-\pi}{2} \right) = \frac{5q}{8\pi\epsilon_0 OP^3} \end{cases}$$

B.2 1. La réponse est 1.

La force s'exprime par $\vec{f} = \frac{1}{4\pi\epsilon_0} \frac{q_1 q_2}{AP^2} \vec{u}_{A \rightarrow P}$. Pour calculer la distance, on développe $AP^2 = AO^2 + OP^2 + 2\vec{AO} \cdot \vec{OP} = 4b^2 \sin^2 \frac{\varphi}{2}$ donc pour la norme $f = \frac{1}{4\pi\epsilon_0} \frac{Q^2}{16b^2 \sin^2 \frac{\varphi}{2}}$.

2. La réponse est 3.

A l'équilibre, la somme des forces est nulle soit en projetant sur la direction perpendiculaire à OP $mg \sin \varphi_e = f \sin \alpha$. Or $2\alpha + \varphi = \pi$. On en déduit en utilisant l'expression de f de la question précédente $\sin^3 \frac{\varphi_e}{2} = \frac{1}{4\pi\epsilon_0} \frac{Q^2}{32mgb^2}$.

3. La réponse est 2.

On projette la nullité de la somme des forces parallèlement à OP et on en déduit la tension $T = mg \cos \varphi_e + \frac{1}{4\pi\epsilon_0} \frac{Q^2}{16b^2 \sin^2 \frac{\varphi_e}{2}}$.

4. La réponse est 4.

On déduit de la relation précédente $Q = \sqrt{4\pi\epsilon_0 32mgb^2 \sin^3 \frac{\varphi_e}{2}} = 4,0 \cdot 10^{-7} \text{ C}$.

5. La réponse est 1.

L'énergie potentielle est la somme de la contribution du poids et de la force coulombienne $E_p = \frac{1}{4\pi\epsilon_0} \frac{Q^2}{8b \sin \frac{\varphi}{2}} - mgb \cos \varphi$.

6. La réponse est 1.

Pour déterminer la position d'équilibre, on résout $\frac{dE_p}{d\varphi} = 0$ soit $\sin^3 \frac{\varphi}{2} = \frac{1}{4\pi\epsilon_0} \frac{Q^2}{32mgb^2}$. Pour sa stabilité, il faut étudier le signe de $\frac{d^2 E_p}{d\varphi^2}$. Ici le calcul de la dérivée seconde de l'énergie potentielle

donne $mgb \cos \varphi + \frac{1}{4\pi\epsilon_0} \frac{Q^2 \left(\cos^2 \frac{\varphi}{2} + \frac{1}{2} \sin^2 \frac{\varphi}{2} \right)}{16b \sin^3 \frac{\varphi}{2}} > 0$ donc l'équilibre est stable.

Chapitre 44

A.1 1. On a invariance par translation selon toutes les directions : le champ électrostatique ne dépend donc d'aucune coordonnée d'espace.

Tout plan passant par M est plan de symétrie : le champ électrostatique appartient à tous ces plans. Il est donc nul.

2. Le flux du champ électrostatique à travers toute surface fermée est nul puisque le champ lui-même est nul. Or à l'intérieur de cette surface, la quantité de charges n'est pas nulle. Le théorème de Gauss n'est donc pas valide. Ceci s'explique par la présence de charges à l'infini, ce qui est contraire aux conditions d'application du théorème de Gauss. L'une des hypothèses nécessaires à la démonstration (hors programme et donc non effectuée dans le cours) du théorème de Gauss est la décroissance du champ électrostatique en $\frac{1}{r^2}$, c'est cette hypothèse qui est mise en défaut ici.

A.2 1. Sur une équipotentielle, le potentiel électrostatique est constant en tout point. Comme $\vec{E} = -\vec{\text{grad}}V$, le module du champ électrostatique dépend de la « vitesse » de variation du potentiel et non de la valeur du potentiel : il n'est donc pas constant sur une équipotentielle.

2. Par définition du gradient et du potentiel électrostatique, on a : $dV = -\vec{E} \cdot d\vec{OM}$. Or sur une équipotentielle, $dV = 0$ donc le produit scalaire du champ électrostatique avec le vecteur tangent est nul. Le champ est donc en tout point perpendiculaire aux équipotentielles.

3. Une équipotentielle peut se recouper en des points où le champ électrostatique est nul (Cf. 2).

4. Deux surfaces équipotentielles correspondent à deux valeurs différentes de potentiel : elles ne peuvent donc pas se couper.

A.3 Les lignes de champ correspondent à celles d'un champ électrostatique si la circulation du champ le long de toute courbe fermée est nulle. C'est le cas des lignes de champ (a), (c), (d) et (f) qui sont celles respectivement d'un champ uniforme, du champ créé par une charge ponctuelle, d'un champ ne dépendant que d'une direction $\vec{E} = E(x)\vec{u}_x$ et du champ créé par une sphère uniformément chargée.

Pour les cas (b) et (e), il suffit de calculer la circulation sur un cercle dont le centre est au centre des lignes de champ pour se persuader que la circulation sur ce cercle n'est pas nulle donc que ces lignes de champ ne peuvent pas être celles d'un champ électrostatique.

A.4 1. Les éléments de surface élémentaire $d\vec{S}$ et $d\vec{S}'$ sur deux faces opposées du cube vérifient $d\vec{S} = -d\vec{S}'$. Dans le cas d'un champ uniforme, la norme, la direction et le sens sont identiques en tout point de l'espace. Les flux relatifs à deux faces opposées sont donc opposés : le flux total est nul. Ceci étant valable pour les trois directions, le flux demandé est nul et $Q_{\text{int}} = 0$.

2. Pour les faces parallèles à (xOy) et (xOz) , le champ électrostatique est perpendiculaire à la surface orientée : le produit scalaire de deux vecteurs perpendiculaires étant nul, le flux du champ électrostatique à travers ces faces est nul.

Pour la face $x = 0$, le champ est nul donc le flux l'est également.

Reste la face $x = a$ sur laquelle on a $\vec{E} = Ca\vec{u}_x$ qui a une valeur constante et est perpendiculaire à la surface. Le flux vaut : $\Phi = Ca^2 = Ca^3$ et $Q_{\text{int}} = C\epsilon_0 a^3$.

3. Pour les mêmes raisons qu'à la question précédente, on ne s'occupe que du calcul du flux à travers la face située à $x = a$ sur laquelle on a $\vec{E} = Ca^2\vec{u}_x$. Le flux vaut : $\Phi = Ca^2 a^2 = Ca^4$ et $Q_{\text{int}} = C\epsilon_0 a^4$.

4. Pour les faces parallèles à (xOy) , le champ électrostatique est perpendiculaire à la surface orientée : le flux est nul.

Pour les faces parallèles à (xOz) , seule la composante sur $\vec{u}_y$ n'est pas perpendiculaire à la surface orientée et donne une contribution non nulle au flux. Cependant à x donné, le champ sur la face

$y = 0$ et sur la face $y = a$ a la même valeur et la même orientation tandis que les surfaces orientées sont opposées. Les contributions au flux s'annulent l'une l'autre et le flux total correspondant est nul. Par un raisonnement analogue, le flux total relatif aux faces parallèles à (yOz) est nul. Le flux total est par conséquent nul et $Q_{\text{int}} = 0$.

B.1 1. On applique le principe de superposition :

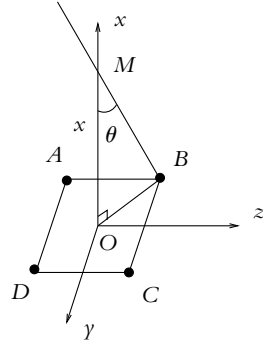
$$V = \frac{q}{4\pi\epsilon_0} \left(\frac{1}{MA} + \frac{1}{MB} + \frac{1}{MC} + \frac{1}{MD} \right).$$

Or $MA = MB = MC = MD = \sqrt{MO^2 + OA^2}$ et

$$OA = \frac{AC}{2} = \frac{a}{\sqrt{2}} \text{ soit}$$

$$MA = \sqrt{x^2 + \frac{a^2}{2}} = \sqrt{\frac{a^2 + 2x^2}{2}}.$$

$$\text{On en déduit : } V = \frac{4q\sqrt{2}}{4\pi\epsilon_0\sqrt{a^2 + 2x^2}}.$$



2. On déduit le champ électrostatique à partir de la relation $\vec{E} = -\vec{\text{grad}}V$. V ne dépend ici que de x , le champ n'aura de composante non nulle que suivant l'axe (Ox) . $\vec{E} = \frac{2q\sqrt{2}}{\pi\epsilon_0} \frac{x}{(a^2 + 2x^2)^{-\frac{3}{2}}} \vec{u}_x$.

3. Les plans AMC et BMD sont des plans de symétrie : le champ électrostatique appartient à l'intersection de ces plans et est donc colinéaire à $\vec{u}_x$. On retrouve une partie du résultat de la question précédente.

4. D'après la question précédente, on ne calcule que la composante suivant l'axe (Ox) donc on projette sur cet axe la relation :

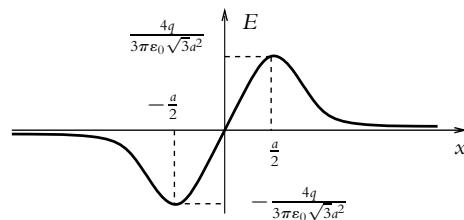
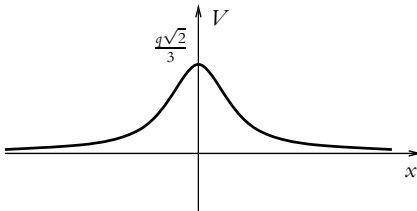
$$\vec{E} = \frac{q}{4\pi\epsilon_0} \left(\frac{\vec{AM}}{AM^3} + \frac{\vec{BM}}{BM^3} + \frac{\vec{CM}}{CM^3} + \frac{\vec{DM}}{DM^3} \right)$$

$$\text{soit } \vec{E} = \frac{q}{4\pi\epsilon_0} \frac{4 \cos \theta}{AM^2} \vec{u}_x = \frac{q}{\pi\epsilon_0} \frac{x}{(a^2 + 2x^2)^{-\frac{3}{2}}} \vec{u}_x.$$

5. Le plan du carré est un plan de symétrie de la distribution de charges donc le potentiel est une fonction paire de x : $V(x) = V(-x)$ et le champ une fonction impaire : $E(x) = -E(-x)$.

6. Cf. figure page suivante.

7. Cf. figure page suivante.



B.2 1. On applique la relation $\vec{E} = -\vec{\text{grad}}V$. V ne dépendant que de r , le champ est radial :

$$\vec{E} = -\frac{dV}{dr} \vec{u}_r = \frac{q}{4\pi\epsilon_0} \left(1 + \frac{r}{a} \right) \frac{\exp\left(-\frac{r}{a}\right)}{r^2} \vec{u}_r.$$

2. $r \gg a$ correspond à la limite $r \rightarrow \infty$ et $\vec{E}$ est alors négligeable.

$r \ll a$ correspond à la limite $r \rightarrow 0$ et $\vec{E} \simeq \frac{q}{4\pi\epsilon_0 r^2} \vec{u}_r$, ce qui est l'expression du champ créé par une charge ponctuelle.

3. $\Phi = 4\pi r^2 E(r) = \frac{q}{4\pi\epsilon_0} \left(1 + \frac{r}{a}\right) \exp\left(-\frac{r}{a}\right)$. En appliquant le théorème de Gauss, on en déduit la charge contenue dans la sphère de rayon r : $Q_{\text{int}} = \frac{q}{4\pi} \left(1 + \frac{r}{a}\right) \exp\left(-\frac{r}{a}\right)$. La distribution de charges ne dépend que de r .

4. Quand r tend vers 0, $q(r)$ tend vers $+q$. Quand r tend vers l'infini, $q(r)$ tend vers 0, la charge $-q$ est donc répartie dans tout l'espace. Elle représente la charge de l'électron.

5. Entre deux sphères de rayon r et $r + dr$, la charge est :

$$dQ_{\text{int}} = Q_{\text{int}}(r + dr) - Q_{\text{int}}(r) = 4\pi r^2 \rho(r) dr$$

soit :

$$\begin{aligned} \rho &= \frac{1}{4\pi r^2} \frac{dQ_{\text{int}}}{dr} = \frac{1}{4\pi r^2} \frac{d}{dr} \left(\frac{q}{4\pi} \left(1 + \frac{r}{a}\right) \exp\left(-\frac{r}{a}\right) \right) \\ &= \frac{q}{r^2} \left(\frac{1}{a} - \frac{1}{a} \left(1 + \frac{r}{a}\right) \right) \exp\left(-\frac{r}{a}\right) = -\frac{q}{4\pi a^2 r} \exp\left(-\frac{r}{a}\right) \end{aligned}$$

6. Puisqu'il s'agit de l'atome d'hydrogène, $q = e$.

D'autre part, $dq(r) = -eP(r) dr = \rho \times 4\pi r^2 dr$ d'où $P(r) = \frac{r}{a^2} \exp\left(-\frac{r}{a}\right)$. La probabilité de trouver l'électron $p(r) = 4\pi r^2 P(r)$. Elle est nulle en $r = 0$ et à l'infini et maximale en $r = a$, "rayon" de l'atome (pour l'atome d'hydrogène, $a = 52$ pm).

B.3 1. Le potentiel créé par une sphère de rayon r uniformément chargée, en un point de sa surface est : $V(r) = \frac{\rho r^2}{3\epsilon_0}$ (Cf. cours). Il faut déplacer la charge $dq = \rho 4\pi r^2 dr$ d'une région de potentiel nul à

une région de potentiel $V(r)$, il faut donc fournir le travail $\delta W = dqV(r) = \rho 4\pi r^2 \frac{\rho r^2}{3\epsilon_0} dr$.

2. En intégrant cette expression entre $r = 0$ et $r = R$ et en utilisant la relation $Q = \rho \frac{4\pi R^2}{3}$, on obtient :

$$E = \frac{3Q^2}{20\pi\epsilon_0 R}.$$

3. L'analogie avec la gravitation s'obtient en remplaçant $\frac{1}{4\pi\epsilon_0}$ par $-G$ et Q par M . On en déduit l'énergie gravitationnelle d'une sphère de rayon R et de masse M : $E = -\frac{3}{5} \frac{GM^2}{R}$.

B.4 1. Tout plan contenant l'axe Oz est plan de symétrie de la distribution. Le champ $\vec{E}$ sur l'axe appartient à ces plans, donc à leur intersection qui est l'axe Oz lui-même.

2. a) Le champ $\vec{E}$ a les mêmes symétries que la distribution de charges, or le plan de la spire est plan de symétrie pour la distribution, donc plan de symétrie pour $\vec{E}$. Ainsi $\vec{E}(-z) = -\vec{E}(z)$.

b) Le champ élémentaire $d\vec{E}$ créé par un élément de spire $d\ell$ autour de P (figure S44.1) est

$$d\vec{E} = \frac{\lambda d\ell}{4\pi\epsilon_0} \frac{\vec{PM}}{PM^3}.$$

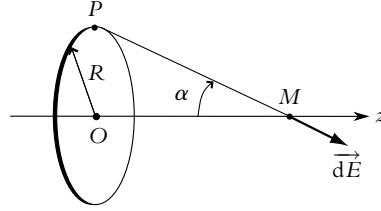


Figure S44.1

On projette sur l'axe Oz et on intègre sur la spire :

$$dE_z = \frac{\lambda}{4\pi\epsilon_0} \frac{\cos \alpha}{PM^2} d\ell \Rightarrow E_z = \frac{\lambda}{4\pi\epsilon_0} \frac{\cos \alpha}{PM^2} \oint_{\text{spire}} d\ell = \frac{\lambda}{2\epsilon_0} \frac{Rz}{(R^2 + z^2)^{3/2}}$$

c) Le graphe de la fonction $E(z)$ est tracé sur la figure S44.2

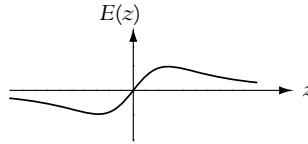


Figure S44.2

3. a) Le plan contenant le point M et l'axe Oz est plan de symétrie de la distribution de charges. Or ce plan a pour vecteurs directeurs les vecteurs $\vec{u}_r$ et $\vec{u}_z$ donc le champ $\vec{E}$ en M qui appartient à ce plan n'a pas de composante suivant $\vec{u}_\theta$.
- b) La distribution de charges est invariante par rotation d'axe Oz donc $\vec{E}$ est indépendant de θ .
- c) Si l'on considère une surface fermée au voisinage de l'axe, elle ne contient pas de charges. Le théorème de Gauss dit que le flux à travers cette surface est nul. Le champ $\vec{E}$ est à circulation conservative.

d) D'après ce que l'on vient d'écrire le flux total φ_t à travers la surface cylindrique est nul puisqu'elle ne contient pas de charges. Le flux se décompose en trois termes : le flux à travers la surface latérale φ_L , celui à travers la base située en z noté $\varphi(z)$ et celui à travers la base située en $z + dz$ noté $\varphi(z + dz)$. On rappelle que les vecteurs surfaces élémentaires sont orientés vers l'extérieur de la surface. Comme les intégrales se font sur des petites surfaces, elles se ramènent au produit de la composante du champ orthogonale à la surface par l'aire de celle-ci.

$$\begin{aligned} \bullet \varphi_L &= \iint_{S_L} \vec{E} \cdot d\vec{S}_L = \iint_{S_L} \vec{E} \cdot dS_L \vec{u}_r = \iint_{S_L} E_r dS_L = E_r(r, z) 2\pi r dz. \\ \bullet \varphi(z) &= \iint_{S_z} \vec{E} \cdot d\vec{S}_z = \iint_{S_z} \vec{E} \cdot (-dS_z \vec{u}_z) = - \iint_{S_z} E_z(z) \cdot dS_z = -\pi r^2 E_z(0, z). \\ \bullet \varphi(z + dz) &= \iint_{S_{z+dz}} \vec{E} \cdot d\vec{S}_{z+dz} = \iint_{S_{z+dz}} \vec{E} \cdot (dS_{z+dz} \vec{u}_z) = \pi r^2 E_z(0, z + dz). \end{aligned}$$

Finalement, on peut écrire le flux total :

$$\begin{aligned} 0 &= E_r(r, z) 2\pi r dz + \pi r^2 (E_z(0, z + dz) - E_z(0, z)) \\ &= E_r(r, z) 2\pi r dz + \pi r^2 \frac{dE_z}{dz}(0, z) \\ \Rightarrow E_r(r, z) &= -\frac{r}{2} \frac{dE_z}{dz}(0, z) \end{aligned}$$

L'expression de E_r au voisinage de l'axe est donc :

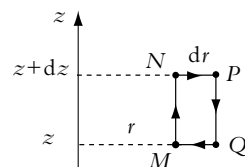
$$E_r(r, z) = \frac{\lambda R r (2z^2 - R^2)}{4\epsilon_0 (R^2 + z^2)^{5/2}}$$

e) La circulation de $\vec{E}$ le long de ce contour est nulle. Elle s'écrit :

$$\mathcal{C} = E_z(r, z) dz + E_r(r, z + dz) dr - E_z(r + dr, z) dz - E_r(r, z) dr = 0$$

On en déduit, après développement limité et simplification par $dr dz$:

$$\frac{\partial E_z}{\partial r}(r, z) = \frac{\partial E_r}{\partial z}(r, z)$$



En utilisant l'expression de $E_r(r, z)$ établie précédemment, on obtient :

$$\frac{\partial E_z}{\partial r}(r, z) = -\frac{r}{2} \frac{d^2 E_z}{dz^2}(0, z)$$

qui s'intègre (entre 0 et r) en :

$$E_z(r, z) = E_z(0, z) - \frac{r^2}{4} \frac{d^2 E_z}{dz^2}(0, z)$$

► **Remarque** : Le champ créé par la spire s'écrit donc :

Sur l'axe : $\vec{E}(M) = E_z(0, z)\vec{u}_z = f(z)\vec{u}_z$.

Au voisinage de l'axe :

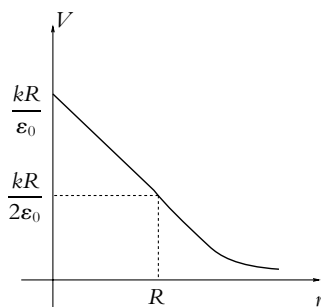
- au premier ordre en r : $\vec{E}(M) = -\frac{r}{2}f'(z)\vec{u}_r + f(z)\vec{u}_z$
- au deuxième ordre en r : $\vec{E}(M) = -\frac{r}{2}f'(z)\vec{u}_r + \left(f(z) - \frac{r^2}{4}f''(z)\right)\vec{u}_z$.

B.5 1. On a invariance par toute rotation autour de O donc le champ $\vec{E}$ ne dépend que de r en coordonnées sphériques.

2. Tout plan passant par O et M est plan de symétrie donc le champ $\vec{E}$ est suivant $\vec{u}_r$.

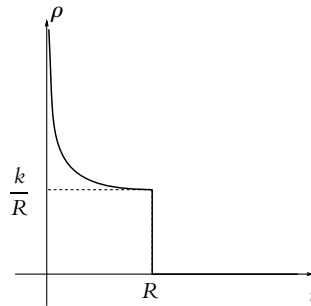
3. D'après la relation $\vec{E} = -\text{grad}V$, on tire $dV = -\vec{E} \cdot d\vec{OM} = E dr$. Par intégration sur r , on en déduit $V = \frac{kR^2}{2\epsilon_0 r}$ pour $r > R$ (la constante d'intégration est nulle car le potentiel V est choisi nul pour r infini) et $V = \frac{kR}{\epsilon_0} \left(1 - \frac{r}{2}\right)$ pour $r < R$ (la constante d'intégration est déterminée par continuité du potentiel en $r = R$). On remarquera que E étant continu en $r = R$, le potentiel ne présente pas de point anguleux en ce point-là.

4.

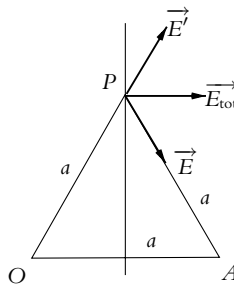


5. On applique la relation fournie dans l'énoncé et on obtient $\rho = \frac{k}{r}$ pour $r < R$ et nul pour $r > R$.

6.



7. On a $dq = 4\pi r^2 \rho dr$ en calculant le volume de la couche sphérique.
8. On intègre la relation précédente entre O et R . On obtient $q_0 = 2\pi k R^2$.
9. L'application du théorème de Gauss à une sphère de rayon $r > R$ et de centre O permet de retrouver le champ créé par une charge ponctuelle. On en déduit le résultat pour une charge q_0 placée en O .
10. La force issue de la loi de Coulomb s'écrit $\vec{f} = \frac{1}{4\pi\epsilon_0} \frac{qq_0}{a^2} \vec{u}$ en utilisant les notations de l'énoncé.
11. Le point P est dans le plan médiateur qui est un plan d'antisymétrie donc le champ est perpendiculaire en P au plan médiateur : il est donc suivant $\vec{u}$.
- 12.



13. Le champ total s'écrit $E_{\text{tot}} = 2E \cos 60^\circ = \frac{1}{4\pi\epsilon_0} \frac{q_0}{a^2}$.
14. Le potentiel total s'écrit $V_{\text{tot}} = \frac{1}{4\pi\epsilon_0} \frac{q_0}{a} + \frac{1}{4\pi\epsilon_0} \frac{-q_0}{a} = 0$.

Chapitre 45

- A.1** 1. Le plan de la feuille est plan de symétrie de la distribution ainsi que le plan perpendiculaire à la feuille et contenant OM . Le champ $\vec{E}$ est suivant $\vec{u}_y$ (figure S45.1).

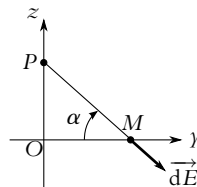


Figure S45.1

Le champ élémentaire créé en M par l'élément de fil dz autour de P est :

$$d\vec{E} = \frac{1}{4\pi\epsilon_0} \frac{\lambda dz \vec{PM}}{PM^3}$$

On projette sur la direction finale du champ Oy soit :

$$dE_y = \frac{1}{4\pi\epsilon_0} \frac{\lambda dz \cos \alpha}{PM^2}$$

or $\tan \alpha = z/D$ soit $dz = \frac{D}{\cos^2 \alpha} d\alpha$. On en déduit :

$$dE_y = \frac{\lambda}{4\pi\epsilon_0 D} \cos \alpha d\alpha \Rightarrow E_y = \int_{-\alpha_1}^{\alpha_1} dE_y = \frac{\lambda}{2\pi\epsilon_0 D} \sin \alpha_1$$

avec $\sin \alpha_1 = L/\sqrt{4D^2 + L^2}$.

2. Les deux plans contenant les diagonales du carré et l'axe Oz sont plans de symétries de la distribution donc $\vec{E}$ appartient à ces plans et aussi à leur intersection : $\vec{E}$ est selon Oz .

Sur la figure S45.2 on a représenté l'un des fils perpendiculaire au plan de la feuille et le champ qu'il crée au point M .

On doit projeter ce champ sur la direction finale et le multiplier par 4 pour tenir compte des quatre fils. Le champ final est donc :

$$E_z = \frac{2\lambda}{\pi\epsilon_0 D} \frac{L}{\sqrt{4D^2 + L^2}} \cos \beta = \frac{2\lambda}{\pi\epsilon_0 D} \frac{L}{\sqrt{4D^2 + L^2}} \frac{z}{D}$$

Dans le cas du carré de côté L , la distance D est égale à

$$\sqrt{z^2 + \left(\frac{L}{2}\right)^2}.$$

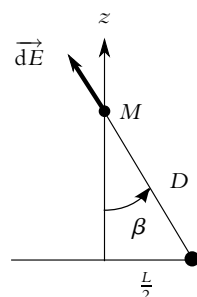


Figure S45.2

- A.2** 1. Tout plan contenant Oz est plan de symétrie de la distribution et de plus le plan perpendiculaire à Oz contenant O (figure S45.3) est plan d'antisymétrie. Le champ en O est selon Oz .

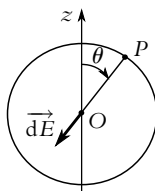


Figure S45.3

2. a) On utilise les coordonnées sphériques et on note R le rayon de la sphère. L'expression du champ élémentaire est :

$$d\vec{E} = \frac{1}{4\pi\epsilon_0} \frac{\sigma dS \vec{PO}}{PO^3}$$

D'après la figure (S45.3), la projection est :

$$dE_z = -\frac{1}{4\pi\epsilon_0} \frac{\sigma dS \cos \theta}{R^2}$$

avec $dS = R^2 \sin \theta d\theta d\varphi$.

- b) L'intégration de l'expression précédente donne :

$$E_z = -\frac{\sigma_0}{4\pi\epsilon_0} \int_{\theta=0}^{\pi} \int_{\varphi=0}^{2\pi} \cos^2 \theta \sin \theta d\theta d\varphi = -\frac{\sigma_0}{4\pi\epsilon_0} \int_0^{\pi} \cos^2 \theta \sin \theta d\theta \int_0^{2\pi} d\varphi = -\frac{\sigma_0}{3\epsilon_0}$$

3. Le point O appartient à un plan d'antisymétrie de la distribution donc $V(O) = 0$.

A.3 Cet exercice redémontre des résultats vus dans le cours. Le corrigé rappelle simplement les résultats que l'on doit obtenir.

1. Pour un fil infini : $\vec{E} = \frac{\lambda}{2\pi r \epsilon_0} \vec{u}_r$ et $V(r) = -\frac{\lambda}{2\pi \epsilon_0} \ln \frac{r}{R_0} + V_0$ où V_0 est le potentiel en $r = R_0$.
2. Pour un plan infini : $\vec{E} = \frac{\sigma}{2\epsilon_0} \vec{u}_z$ si $z > 0$ et $\vec{E} = -\frac{\sigma}{2\epsilon_0} \vec{u}_z$ si $z < 0$.

Le calcul du potentiel donne pour $z > 0$:

$$\frac{dV}{dz} = -\frac{\sigma}{2\epsilon_0} \Rightarrow V(z) = -\frac{\sigma z}{2\epsilon_0} + V_0$$

Si l'on note V_0 le potentiel en $z = 0$ et pour $z < 0$:

$$\frac{dV}{dz} = \frac{\sigma}{2\epsilon_0} \Rightarrow V(z) = \frac{\sigma z}{2\epsilon_0} + V_0$$

3. Pour une boule uniformément chargée en volume :

- pour $r < R$, $\vec{E} = \frac{\rho}{3\epsilon_0} r \vec{u}_r$;
- pour $r > R$, $\vec{E} = \frac{\rho}{3\epsilon_0} \frac{R^3}{r^2} \vec{u}_r$.

Le calcul du potentiel donne :

- pour $r > R$, $\frac{dV}{dr} = -\frac{\rho}{3\epsilon_0 r^2} R^3 \Rightarrow V = \frac{\rho}{3\epsilon_0 r} R^3 + cte$, si on prend le potentiel nul à l'infini, la constante est nulle ;
- pour $r < R$, $\frac{dV}{dr} = -\frac{\rho}{3\epsilon_0} r \Rightarrow V = -\frac{\rho}{6\epsilon_0} r^2 + cte$, on détermine la constante par continuité en $r = R$ soit :

$$\frac{R^2 \rho}{3\epsilon_0} = \frac{-\rho R^2}{6\epsilon_0} + cte \Rightarrow cte = \frac{R^2 \rho}{2\epsilon_0}$$

A.4 Le système est représenté sur la figure (S45.4). Il ne possède pas de fortes symétries.

On va utiliser le théorème de superposition en décomposant le système en une boule pleine $\mathcal{B}_1$ de rayon R , de centre O , de densité volumique ρ et une boule pleine $\mathcal{B}_2$ de rayon a , de centre O' , de densité volumique $-\rho$. À l'endroit de la petite boule, on aura bien par superposition une charge nulle ($-\rho + \rho$).

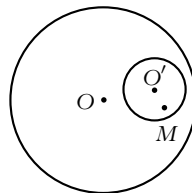


Figure S45.4

D'autre part, on utilise le résultat du calcul du champ créé par une boule (voir cours). Un point M intérieur à la cavité est intérieur à $\mathcal{B}_1$ et à $\mathcal{B}_2$. Le champ créé par $\mathcal{B}_1$ en M est $\vec{E}_1 = \frac{\rho}{3\epsilon_0} \vec{OM}$ et le champ créé par $\mathcal{B}_2$ est $\vec{E}_2 = \frac{-\rho}{3\epsilon_0} \vec{O'M}$. Le champ total est donc :

$$\vec{E} = \vec{E}_1 + \vec{E}_2 = \frac{\rho}{3\epsilon_0} (\vec{OM} - \vec{O'M}) = \frac{\rho}{3\epsilon_0} \vec{OO'}$$

Le champ est uniforme dans la cavité.

- A.5** 1. La distribution est à symétrie sphérique, donc on prend comme volume élémentaire, le volume compris entre une sphère de rayon r et une sphère de rayon $r + dr$. Le volume considéré est égal à : $dV(r) = 4\pi r^2 dr$ (aire de la sphère de rayon r multipliée par l'épaisseur dr). La charge totale est donc :

$$Q = \int_0^a \rho(r) dV(r) = 4\pi \rho_0 \int_0^a \left(1 - \frac{r^2}{a^2}\right) r^2 dr = \frac{8\pi \rho_0 a^3}{15}$$

2. Pour tout point P de l'espace, tout plan contenant la droite OP est plan de symétrie de la distribution de charge. Le champ $\vec{E}$ en P appartient à tous ces plans donc à leur intersection. Le champ en P est radial (selon $\vec{u}_r$ en coordonnées sphériques).

La distribution est invariante par toute rotation autour de O , donc $\vec{E}$ est indépendant de θ et φ , il ne dépend que de r . Ainsi $\vec{E}(P) = E(r)\vec{u}_r$.

3. On utilise le théorème de Gauss. Étant données les symétries et invariance de $\vec{E}$, les surfaces de Gauss choisies sont des sphères de centre O et de rayon r . Le flux à travers cette surface est :

$$\varphi = \oiint \vec{E}(P) \cdot d\vec{S}_P = \oiint E(r)\vec{u}_r \cdot dS_P \vec{u}_r = \oiint E(r) dS_P = 4\pi r^2 E(r)$$

Il faut maintenant calculer la charge intérieure à cette sphère. Si $r > a$, la charge est celle calculée précédemment et on trouve :

$$\vec{E} = \frac{Q}{4\pi\epsilon_0 r^2} \vec{u}_r$$

4. Pour $r < a$ il faut calculer la charge intérieure par l'intégrale :

$$Q(r) = \int_0^r \rho(r) dV(r) = 4\pi\rho_0 \int_0^r \left(1 - \frac{r^2}{a^2}\right) r^2 dr = \frac{4\pi\rho_0 r^3}{3} \left(1 - \frac{3r^2}{5a^2}\right)$$

On obtient donc :

$$4\pi r^2 E(r) = \frac{Q(r)}{\epsilon_0} \Rightarrow \vec{E}(r) = \frac{\rho_0 r}{3\epsilon_0} \left(1 - \frac{3r^2}{5a^2}\right) \vec{u}_r$$

5. et 6. Pour calculer le potentiel, on utilise la relation $\vec{E} = -\vec{\text{grad}} V$ ce qui donne ici $\frac{dV}{dr} = -E$. Pour $r > a$ on retrouve le potentiel d'une charge ponctuelle $V = Q/(4\pi\epsilon_0 r)$. Pour $r < a$, il vient :

$$\frac{dV}{dr} = -\frac{\rho_0 r}{3\epsilon_0} \left(1 - \frac{3r^2}{5a^2}\right) \Rightarrow V(r) = \frac{\rho_0}{\epsilon_0} \left(-\frac{r^2}{6} + \frac{r^4}{20a^2}\right) + cte$$

On détermine la constante par continuité du potentiel en $r = a$, soit :

$$-\frac{7a^2\rho_0}{60\epsilon_0} + cte = \frac{2\rho_0 a^2}{15\epsilon_0} \Rightarrow cte = \frac{\rho_0 a^2}{4\epsilon_0}$$

- A.6** 1. Le problème est à symétrie sphérique. Le champ électrostatique à l'intérieur de la boule est $\vec{E} = E_0 \vec{u}_r$. Le théorème de Gauss à travers une sphère de centre O et de rayon $r < R$ donne la charge $q(r)$ contenue dans cette sphère : $q(r) = 4\pi\epsilon_0 r^2 E_0$. Or : $q(r) = \int_0^r \rho(u) 4\pi u^2 du$. En dérivant cette équation par rapport à r , on obtient : $\rho(r) = \frac{2\epsilon_0 E_0}{r}$.

2. Pour $r < R$: $\vec{E} = E_0 \vec{u}_r$; pour $r > R$ on utilise le théorème de Gauss et on trouve : $\vec{E} = E_0 R^2 / r^2 \vec{u}_r$. Pour déterminer v , on utilise la relation $\vec{E} = -\vec{\text{grad}} V$ avec V continu en $r = R$ et V tend vers 0 loin de la boule. On en déduit : pour $r < R$, $V(r) = -E_0 r$ et pour $r > R$, $V(r) = \frac{E_0 R^2}{r^2}$.

- A.7** 1. Le plan xOz est un plan de symétrie de la distribution de charges ainsi que le plan yOz (distribution infinie). Le champ $\vec{E}$ en M appartient à ces deux plans donc à leur intersection : $\vec{E}$ est suivant Oz .

2. a) La charge de la distribution élémentaire doit être la même pour les deux modélisations donc $\lambda = \sigma dy$.
b) On utilise l'expression du champ créé par un fil démontrée en cours. Avec les notations de la figure (S45.5), on peut écrire :

$$d\vec{E} = \frac{\lambda}{2\pi\epsilon_0 PM^2} \vec{PM} \Rightarrow dE_z = \frac{\sigma \cos \theta dy}{2\pi\epsilon_0 PM} = \frac{z\sigma}{2\pi\epsilon_0} \frac{dy}{y^2 + z^2}$$

c) On intègre l'expression de dE_z :

$$\begin{aligned} E_z &= \frac{z\sigma}{2\pi\epsilon_0} \int_{-b/2}^{b/2} \frac{dy}{y^2 + z^2} \\ &= \frac{z\sigma}{2\pi\epsilon_0} \frac{1}{z} \left[\arctan \frac{y}{z} \right]_{-b/2}^{b/2} \\ &= \frac{\sigma}{\pi\epsilon_0} \arctan \frac{b}{2z} \end{aligned}$$

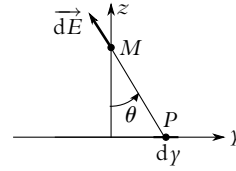


Figure S45.5

A.8 1. On a représenté la distribution en coupe dans le plan yOz sur la figure (S45.6).

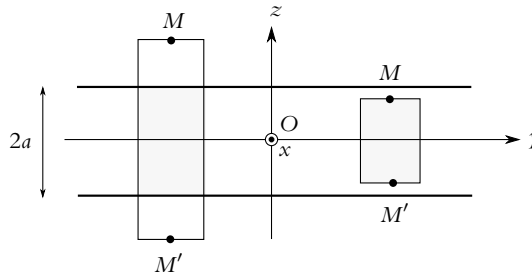


Figure S45.6

Les symétries et les invariances sont les mêmes que pour le plan infini (voir cours). Le champ $\vec{E}$ est suivant $\vec{u}_z$ et ne dépend que de z . D'autre part, comme le plan xOy est plan de symétrie, on a $E(-z) = E(z)$.

On choisit comme surface de Gauss des parallélépipèdes symétriques comme dans le cas du plan infini. Le calcul du flux est le même et on trouve $\Phi = 2ES$ où E est le champ en M de cote z positive et S la surface de la base des parallélépipèdes.

Le calcul de la charge intérieure diffère en revanche de celle du plan infini. On a représenté en coupe deux cas de parallélépipèdes : celui de droite a une hauteur $2z < 2a$ et celui de gauche une hauteur $2z > 2a$. Dans les deux cas, on a grisé la zone correspondant à la charge intérieure.

Pour $z < a$, on a $Q_{\text{int}} = 2zS\rho$ où ρ est la densité volumique de charges. Le théorème de Gauss donne alors $2ES = 2zS\rho/\epsilon_0$ soit $E = z\rho/\epsilon_0$.

Pour $z > a$, on a $Q_{\text{int}} = 2aS\rho$. Le théorème de Gauss donne alors $2ES = 2aS\rho/\epsilon_0$ soit $E = a\rho/\epsilon_0$.

On obtient le champ pour $z < 0$ par symétrie. On récapitule donc :

- Si $z > a$, $\vec{E} = \frac{\rho a}{\epsilon_0} \vec{u}_z$.
- Si $a > z > -a$, $\vec{E} = \frac{\rho z}{\epsilon_0} \vec{u}_z$ (l'orientation est donnée par le signe de z).
- Si $z < -a$, $\vec{E} = -\frac{\rho a}{\epsilon_0} \vec{u}_z$.

2. Le calcul du potentiel se fait avec $\vec{E} = -\vec{\text{grad}} V$. On fixe par exemple le potentiel nul en $z = 0$. On commence par la zone $|z| < a$:

$$\frac{dV}{dz} = -\frac{\rho z}{\epsilon_0} \Rightarrow V = -\frac{\rho z^2}{2\epsilon_0} + cte$$

La constante est nulle si on prend $V(0) = 0$.

On continue avec $z > a$:

$$\frac{dV}{dz} = -\frac{\rho a}{\epsilon_0} \Rightarrow V = -\frac{\rho a z}{\epsilon_0} + cte$$

On détermine la constante par continuité en $z = a$ ce qui donne $cte = \frac{\rho a^2}{2\epsilon_0}$.

Enfin pour $z < -a$:

$$\frac{dV}{dz} = \frac{\rho a}{\epsilon_0} \Rightarrow V = \frac{\rho a z}{\epsilon_0} + cte$$

On détermine la constante par continuité en $z = -a$ ce qui donne aussi $cte = \frac{\rho a^2}{2\epsilon_0}$.

A.9 1. La distribution est invariante par translation le long des axes (Oy) et (Oz) : le potentiel et le champ électrostatiques ne dépendent donc que de x : $V = V(x)$ et $\vec{E} = \vec{E}(x)$.

Tout plan dont l'une des directions est l'axe (Ox) est plan de symétrie de la distribution : le champ est donc dirigé selon (Ox) soit $\vec{E} = E(x)\vec{u}_x$.

Le plan (yOz) est un plan d'antisymétrie de la distribution donc $E(x) = E(-x)$. On se limite donc dans la suite au cas $x \geq 0$.

2. Pour calculer le champ, on applique le principe de superposition en considérant les deux distributions : $\rho_1 = \rho_0$ si $0 \leq x \leq a$ et $\rho_2 = -\rho_0$ si $-a \leq x \leq 0$ et on utilise les résultats de l'exercice précédent. Pour la distribution ρ_1 :

$$E_1\left(\frac{a}{2} + d\right) = \begin{cases} \frac{\rho_0 d}{\epsilon_0} & \text{si } |d| \leq \frac{a}{2} \\ \frac{\rho_0 a}{2\epsilon_0} & \text{si } |d| > \frac{a}{2} \end{cases} \quad \text{soit} \quad E_1(x) = \begin{cases} \frac{\rho_0}{\epsilon_0} \left(x - \frac{a}{2}\right) & \text{si } 0 \leq x \leq a \\ \frac{\rho_0 a}{2\epsilon_0} & \text{sinon} \end{cases}$$

Pour la distribution ρ_2 :

$$E_2(x) = \begin{cases} \frac{-\rho_0}{\epsilon_0} \left(x + \frac{a}{2}\right) & \text{si } -a \leq x \leq 0 \\ \frac{-\rho_0 a}{2\epsilon_0} & \text{sinon} \end{cases}$$

Finalement, on en déduit :

$$E(x) = E_1(x) + E_2(x) = \begin{cases} \frac{\rho_0}{\epsilon_0} (|x| - a) & \text{si } -a \leq x \leq a \\ 0 & \text{si } |x| > a \end{cases}$$

3. Pour calculer le potentiel, on utilise la relation $\vec{E} = -\vec{\text{grad}}V$ donc

$$V = V_0 + \int_0^x E(x) dx = \begin{cases} V_0 + \frac{\rho_0}{\epsilon_0} \frac{x^2 - 2ax}{2} & \text{si } 0 \leq x \leq a \\ V_0 - \frac{\rho_0 a^2}{2\epsilon_0} & \text{si } x \geq a \\ V_0 - \frac{\rho_0}{\epsilon_0} \frac{x^2 + 2ax}{2} & \text{si } -a \leq x \leq 0 \\ V_0 - \frac{3\rho_0 a^2}{2\epsilon_0} & \text{si } x \leq -a \end{cases}$$

Remarque : on a dû donner *a priori* la valeur du potentiel en un point ; en effet, il y a des charges à l'infini donc le potentiel V ne tend pas vers 0 quand x tend vers l'infini.

B.1 1. Le calcul du champ créé par un plan infini a été fait dans le cours. Il vaut $\frac{\sigma}{2\epsilon_0} \vec{u}_z$ pour $z > h$ et $-\frac{\sigma}{2\epsilon_0} \vec{u}_z$ pour $z < h$.

2. On utilise le théorème de superposition. Le champ créé par le plan (chargé $-\sigma$) représentant le sol est $\vec{E}_s = -\frac{\sigma}{2\epsilon_0} \vec{u}_z$. Le champ total est la somme des deux, soit pour $0 < z < h$:

$$\vec{E} = -\frac{\sigma}{\epsilon_0} \vec{u}_z$$

3. La relation $\vec{E} = -\vec{\text{grad}} V$ donne ici :

$$-\frac{dV}{dz} = -\frac{\sigma}{\epsilon_0} \Rightarrow V(z) = \frac{\sigma z}{\epsilon_0} + \text{cte}$$

Sachant qu'on pose $V(0) = 0$, on trouve que la constante est nulle.

4. La différence de potentiel est $\Delta V = V(h) - V(0) = \sigma h / \epsilon_0$. La charge portée par une armature du condensateur est $Q = \sigma S$ où S est la surface du nuage. On en déduit la capacité :

$$C = \frac{Q}{\Delta V} = \frac{S\epsilon_0}{h} = 0,44 \mu\text{F}$$

5. Avec les résultats précédents, on remarque que $\Delta V = V_0 = |E|h$. L'application numérique donne $V_0 = 50.10^6 \text{ V}$.

- B.2** 1. On a invariance par translation le long de l'axe du cylindre et par rotation autour de cet axe donc le champ et le potentiel ne sont fonction que de r en coordonnées cylindriques.

Le plan perpendiculaire à l'axe et le plan contenant l'axe sont des plans de symétrie donc le champ est radial.

On applique le théorème de Gauss à un cylindre d'axe l'axe du conducteur, de rayon r et de hauteur h soit

$$2\pi rhE(r) = \frac{Q_{\text{int}}}{\epsilon_0}$$

avec $Q_{\text{int}} = \pi r^2 h \rho$ pour $r < a$ et $Q_{\text{int}} = \pi a^2 h \rho$ pour $r > a$.

On en déduit $\vec{E}(r) = \frac{\rho r}{2\epsilon_0} \vec{u}_r$ pour $r < a$ et $\vec{E}(r) = \frac{\rho a^2}{2\epsilon_0 r} \vec{u}_r$ pour $r > a$.

2. On obtient le potentiel en intégrant $dV = -E(r)dr$. En utilisant $V(a) = 0$ pour déterminer les constantes, on a $V = \frac{\rho}{4\epsilon_0} (a^2 - r^2)$ pour $r < a$ et $V = -\frac{\rho a^2}{2\epsilon_0} \ln \frac{r}{a}$ pour $r > a$.

3. La différence de potentiel vaut $\Delta U = -\frac{\rho a^2}{2\epsilon_0} \ln \frac{r}{a} = 7,8 \text{ mV}$.

4. $V_1 = -\frac{\rho a^2}{2\epsilon_0} \ln \frac{2(D-x)}{D}$ et $V_2 = \frac{\rho a^2}{2\epsilon_0} \ln \frac{2x}{D}$. On en déduit $V = \frac{\rho a^2}{2\epsilon_0} \ln \frac{x}{D-x}$.

5. Les variations de $f(x) = \ln \frac{x}{D-x}$ pour x entre 0 et D sont obtenues par le calcul de la dérivée $f'(x) = \frac{D}{x(D-x)}$. Cette dérivée est positive entre 0 et D donc la fonction f est croissante de $-\infty$ à $+\infty$ avec $f(x) = 0$ pour $x = \frac{D}{2}$.

- B.3** 1. On a invariance par rotation autour de l'axe Oz donc on utilise les coordonnées cylindriques et on a indépendance par rapport à θ .

2. Il n'y a pas de plan d'antisymétrie. Tout plan contenant l'axe Oz est plan de symétrie. Par conséquent, pour tout point de l'axe Oz , le champ est suivant l'axe Oz . Sinon en un point quelconque, le champ est contenu dans le plan OMz .

3. D'après les résultats de la question précédente, le champ est suivant l'axe Oz .

4. On pose le calcul de l'intégrale $\vec{E}(O) = \frac{1}{4\pi\epsilon_0} \iiint_{P \in V} \frac{\rho d\tau_P \vec{PO}}{PO^3}$. On projette sur l'axe Oz en utilisant le fait que $d\tau_P = r^2 \sin \theta dr d\theta d\varphi$. Le calcul s'effectue en intégrant r entre 0 et R , θ entre 0 et $\frac{\pi}{2}$ et φ entre 0 et 2π . On obtient $\vec{E}(O) = -\frac{\rho R}{4\epsilon_0} \vec{u}_z$.

5. On a les mêmes symétries et les mêmes invariances que précédemment; les conclusions restent donc les mêmes. La surface élémentaire étant donnée par $r dr d\theta$, l'intégration sur le disque donne

$$\vec{E} = -\frac{\sigma}{2\epsilon_0} \left(\frac{z}{\sqrt{z^2 + R^2}} - 1 \right) \vec{u}_z.$$

6. On a ici $\sigma = \rho dz$. Par ailleurs, R dépend de z par la relation $z^2 + R(z)^2 = R^2$ donc $dE = \frac{\rho dz}{2\epsilon_0} \left(\frac{z}{R} - 1 \right)$. L'intégration sur z entre 0 et r donne le résultat.
7. Les quantités $\vec{E}$, q , ρ et $\frac{1}{4\pi\epsilon_0}$ sont les analogues respectivement de $\vec{g}$, m , ρ et $-G$.
8. Le champ $\vec{g}_0$ s'obtient par superposition de $\vec{g}_1$ et $\vec{g}'$ défini comme le champ créé par une demi-sphère. On utilise l'analogie pour écrire $g' = \pi GR (\rho_0 - \rho_1)$ et $\Delta g = -g'$.
9. L'application numérique donne $\Delta g \simeq 4.10^{-4} \text{ m.s}^{-2}$.

B.4 1. a) Tout plan contenant l'axe Oz est plan de symétrie de la distribution de charges, donc le champ $\vec{E}$ appartient à l'axe Oz .

Le champ élémentaire $d\vec{E}$ en M créé par la surface élémentaire dS_p autour de P est $d\vec{E} = \frac{\sigma dS_p}{4\pi\epsilon_0 PM^3} \vec{PM}$. On projette ce champ sur la direction finale Oz :

$$dE_z = \frac{\sigma dS_p \cos \alpha}{4\pi\epsilon_0 PM^2} = \frac{\sigma}{4\pi\epsilon_0} \frac{zR dR d\theta}{(R^2 + z^2)^{3/2}}$$

et on intègre sur le disque :

$$E_z = \frac{z\sigma}{4\pi\epsilon_0} \int_0^{2\pi} d\theta \int_0^r \frac{R dR}{(R^2 + z^2)^{3/2}} = \frac{z\sigma}{2\epsilon_0} \left(\frac{1}{|z|} - \frac{1}{\sqrt{r^2 + z^2}} \right)$$

b) Pour le calcul du potentiel, il vient :

$$dV = \frac{\sigma dS_p}{4\pi\epsilon_0 PM} \Rightarrow V = \frac{\sigma}{4\pi\epsilon_0} \int_0^{2\pi} d\theta \int_0^r \frac{R dR}{\sqrt{R^2 + z^2}} = \frac{\sigma}{2\epsilon_0} (\sqrt{r^2 + z^2} - |z|)$$

c) Dans les expressions précédentes, il faut remplacer σ par ρdz .

2. a) On décompose le cône en disque d'épaisseur dz (figure S45.8).

La variable z des expressions précédentes est égale ici à $O'S$, il faut donc la remplacer par $h - z$ si l'on prend l'origine en O . Si l'on note r le rayon de chaque disque, on peut écrire $\tan \alpha = r/(h - z)$. L'expression du champ créé par un disque d'épaisseur dz est donc :

$$\begin{aligned} dE_z &= \frac{\rho dz}{2\epsilon_0} \left(\frac{h-z}{|h-z|} - \frac{h-z}{\sqrt{r^2 + (h-z)^2}} \right) \\ &= \frac{\rho dz}{2\epsilon_0} \left(1 - \frac{h-z}{\sqrt{(h-z)^2 (1 + \tan^2 \alpha)}} \right) \end{aligned}$$

On peut encore simplifier et l'on obtient :

$$dE_z = \frac{\rho dz}{2\epsilon_0} \left(1 - \frac{1}{\sqrt{1 + \tan^2 \alpha}} \right) = \frac{\rho dz}{2\epsilon_0} (1 - \cos \alpha)$$

L'intégrale est facile à calculer :

$$E_z = \frac{\rho}{2\epsilon_0} (1 - \cos \alpha) \int_0^h dz = \frac{\rho h}{2\epsilon_0} (1 - \cos \alpha)$$

b) Pour le potentiel, il vient :

$$dV = \frac{\rho dz}{2\epsilon_0} \left(\sqrt{(h-z)^2 (1 + \tan^2 \alpha)} - (h-z) \right) = \frac{\rho dz}{2\epsilon_0} (h-z) \left(\sqrt{1 + \tan^2 \alpha} - 1 \right)$$

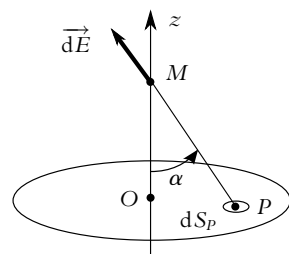


Figure S45.7

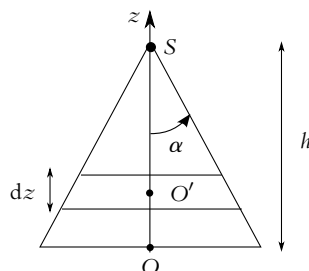


Figure S45.8

et

$$V = \frac{\rho}{2\epsilon_0} \left(\sqrt{1 + \tan^2 \alpha} - 1 \right) \int_0^h (h - z) dz = \frac{h^2 \rho}{4\epsilon_0} \left(\sqrt{1 + \tan^2 \alpha} - 1 \right)$$

3. a) La force de Coulomb entre deux charges est $\vec{f}_c = \frac{q_1 q_2}{4\pi\epsilon_0} \frac{\vec{M}_1 \vec{M}_2}{(M_1 M_2)^3}$ et celle de Newton entre deux

masses $\vec{f}_n = -Gm_1 m_2 \frac{\vec{M}_1 \vec{M}_2}{(M_1 M_2)^3}$. On passe de l'une à l'autre en remplaçant la charge par la masse et $1/(4\pi\epsilon_0)$ par $-G$. Le champ gravitationnel créé en M par une masse ponctuelle m situé en O est $-Gm \frac{\vec{OM}}{(OM)^3}$.

- b) Le champ gravitationnel créé par la Terre est donc $\vec{g} = -\frac{GM_T}{r^2} \vec{u}_r$.

- c) On utilise l'expression du champ électrique créé par un cône où l'on remplace ρ par μ la masse volumique et $1/\epsilon_0$ par $-4\pi G$, ce qui donne $\vec{g}_v = -2\pi G\mu h(1 - \cos \alpha) \vec{u}_z$. Localement le champ créé par la Terre est suivant $\vec{u}_z$ donc le champ total au sommet du volcan est :

$$\vec{g}_t = \left(-2\pi G\mu h(1 - \cos \alpha) - \frac{GM_T}{(R_T + h)^2} \right) \vec{u}_z$$

- d) Le champ de pesanteur sans volcan est $9,760 \text{ m.s}^{-2}$ et avec volcan $9,766 \text{ m.s}^{-2}$. La présence du volcan n'arrive pas à compenser la diminution de g due à l'altitude.

B.5 1. Voir cours.

2. a) Puisque la jonction est globalement neutre, la charge totale est nulle soit (par unité de surface) :

$$\rho_1 L_1 + \rho_2 L_2 = 0$$

- b) Il faut remplacer σ par ρdz dans les expressions du champ créé par un plan mince. Donc si l'axe Oz est perpendiculaire au plan, on a :

- Pour un point au-dessus du plan d $\vec{E} = \frac{\rho dz}{2\epsilon_0 \epsilon_r} \vec{u}_z$.
- Pour un point au-dessous du plan d $\vec{E} = -\frac{\rho dz}{2\epsilon_0 \epsilon_r} \vec{u}_z$.

- c) On a donc 4 zones :

A. Pour $z < -L_1$, le point où l'on calcule le champ est en dessous des plans chargés soit :

$$\vec{E} = \frac{-1}{2\epsilon_0 \epsilon_r} \left(\int_{-L_1}^0 \rho_1 dz + \int_0^{L_2} \rho_2 dz \right) \vec{u}_z = \frac{-1}{2\epsilon_0 \epsilon_r} (\rho_1 L_1 + \rho_2 L_2) \vec{u}_z = \vec{0}$$

B. Pour $-L_1 < z < 0$, il faut partager la première intégrale car M est en-dessous d'une partie de la distribution ρ_1 et au-dessus de l'autre partie. En revanche, le point M est en-dessous de la distribution ρ_2 :

$$\vec{E} = \frac{1}{2\epsilon_0 \epsilon_r} \left(\int_{-L_1}^z \rho_1 dz - \int_z^0 \rho_1 dz - \int_0^{L_2} \rho_2 dz \right) \vec{u}_z = \frac{(\rho_1(2z + L_1) - \rho_2 L_2)}{2\epsilon_0 \epsilon_r} \vec{u}_z$$

C. Pour $0 < z < L_2$, c'est la deuxième intégrale qu'il faut couper en deux :

$$\vec{E} = \frac{1}{2\epsilon_0 \epsilon_r} \left(\int_{-L_1}^0 \rho_1 dz + \int_0^z \rho_2 dz - \int_z^{L_2} \rho_2 dz \right) \vec{u}_z = \frac{(\rho_1 L_1 + \rho_2(2z - L_2))}{2\epsilon_0 \epsilon_r} \vec{u}_z$$

D. Pour $z > L_2$, le point M est au-dessus de toute la distribution :

$$\vec{E} = \frac{1}{2\epsilon_0 \epsilon_r} \left(\int_{-L_1}^0 \rho_1 dz + \int_0^{L_2} \rho_2 dz \right) \vec{u}_z = \frac{1}{2\epsilon_0 \epsilon_r} (\rho_1 L_1 + \rho_2 L_2) \vec{u}_z = \vec{0}$$

L'allure de E est tracée sur la figure (S45.9).

- d) On utilise la relation $\vec{E} = -\vec{\text{grad}} V$ soit ici $\frac{dV}{dz} = -E$.

Pour déterminer les constantes d'intégration, on utilise la continuité du potentiel en $z = -L_1$, $z = 0$ et $z = L_2$. En utilisant la relation d'électroneutralité pour l'expression des résultats, on trouve :

A. Pour $z < -L_1$, $V = \frac{\rho_1 L_1^2}{2\epsilon_0 \epsilon_r}$

B. Pour $-L_1 < z < 0$, $V = \frac{-\rho_1}{2\epsilon_0 \epsilon_r} (2L_1 z + z^2)$

C. Pour $0 < z < L_2$, $V = \frac{\rho_2}{2\epsilon_0 \epsilon_r} (2L_2 z - z^2)$

D. Pour $z > L_2$, $V = \frac{\rho_2 L_2^2}{2\epsilon_0 \epsilon_r}$

- e) La différence de potentiel est :

$$V_0 = V(L_2) - V(L_1) = \frac{1}{2\epsilon_0 \epsilon_r} (\rho_2 L_2^2 - \rho_1 L_1^2) = \frac{L_1 L_2}{2\epsilon_0 \epsilon_r} (\rho_2 - \rho_1)$$

3. a) Tous les atomes d'antimoine sont ionisés donc $\rho_2 = N_2 e$ et tous les atomes de bore sont aussi ionisés donc $\rho_1 = -N_1 e$.

- b) La largeur de la jonction est $\delta = L_1 + L_2$, or $\rho_1 L_1 + \rho_2 L_2 = 0$ donc $L_1 = -\frac{\rho_2}{\rho_1} L_2 = \frac{N_2}{N_1} L_2$ donc $L_1 \ll L_2$. On en déduit $\delta \simeq L_2$.

- c) L'expression de V_0 donne :

$$V_0 = \frac{L_1 L_2 e}{2\epsilon_0 \epsilon_r} (N_1 + N_2) \simeq \frac{N_2 L_2^2 e}{2\epsilon_0 \epsilon_r} = \frac{N_2 \delta^2 e}{2\epsilon_0 \epsilon_r}$$

où l'on a utilisé la relation $N_1 L_1 = N_2 L_2$. On en déduit δ et l'application numérique donne $\delta = 0,58 \mu\text{m}$.

- B.6** 1. a) Le système est invariant par translation d'axe Oz , donc $\vec{E}$ et V sont indépendants de z .

- b) Tout plan perpendiculaire à Oz est plan de symétrie de la distribution de charges, donc $\vec{E}$ contenu dans les plans de symétrie n'a pas de composantes suivant Oz .

- c) Les trois plans passant par l'axe central et les deux électrodes opposées sont des plans de symétrie de la distribution. Le champ $\vec{E}$ en tout point de ces plans est contenu dans le plan. Avec la conclusion de la question précédente, on peut déduire qu'en tout point de ces trois plans $\vec{E}$ est radial.

- d) Les trois plans passant par l'axe central et à égale distance des électrodes sont plans d'antisymétrie de la distribution de charge. Le champ $\vec{E}$ est donc orthogonal à ces plans. Les surfaces auxquelles $\vec{E}$ est orthogonales sont les équipotentielles, donc ces plans sont équipotentiels.

- e) Si l'on fait tourner le système de $2\pi/3$ il est inchangé. D'autre part V est indépendant de z . La période est donc le tiers de la période du cosinus ou du sinus, on peut donc écrire V sous la forme :

$$V(r, \theta, z) = f(r) \sum_p (A_p \cos 3p\theta + B_p \sin 3p\theta)$$

où p est un entier.

2. Le calcul se ramène à celui d'un fil infini de densité linéique λ . Ce calcul est effectué dans le cours. On trouve :

$$\vec{E} = \frac{\lambda}{2\pi\epsilon_0 D} \vec{u}_r \text{ et } V = -\frac{\lambda}{2\pi\epsilon_0} \ln \frac{D}{D_0} + V_0$$

où $\vec{u}_r$ est le vecteur radial par rapport au fil et D_0 la distance au fil du point de potentiel V_0 .

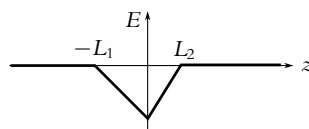


Figure S45.9

3. On choisit le potentiel nul sur l'axe, donc $D_0 = R$ et $V_0 = 0$. En tenant compte de l'alternance des signes des densités linéiques, on a :

$$V(P) = \frac{\lambda}{2\pi\epsilon_0} \left(-\ln \frac{D_1}{R} + \ln \frac{D_2}{R} - \ln \frac{D_3}{R} + \ln \frac{D_4}{R} - \ln \frac{D_5}{R} + \ln \frac{D_6}{R} \right)$$

ce qui donne l'expression demandée en regroupant les logarithmes.

4. On peut écrire les relations suivantes :

$$D_1 = |R - \underline{Z}|; D_3 = |jR - \underline{Z}|; D_5 = |j^2R - \underline{Z}| \Rightarrow D_1 D_3 D_5 = |(R - \underline{Z})(jR - \underline{Z})(j^2R - \underline{Z})|$$

or R, jR et j^2R sont les racines de l'équation $\underline{Z}^3 - R^3 = 0$, donc le produit précédent est proportionnel à $|\underline{Z}^3 - R^3|$. Le terme de plus haut degré est 1 dans les deux cas donc les deux expressions sont égales. Le raisonnement est semblable pour $D_2 D_4 D_6$ avec les racines $-R, -jR$ et $-j^2R$ du polynôme $\underline{Z}^3 + R^3$. On en déduit donc l'expression demandée.

5. En utilisant l'écriture exponentielle de $\underline{Z}$, il vient :

$$\left| \frac{R^3 + \underline{Z}^3}{R^3 - \underline{Z}^3} \right| = \left| \frac{R^3 + r^3 e^{i3\theta}}{R^3 - r^3 e^{i3\theta}} \right| = \left| \left(1 + \frac{r^3}{R^3} e^{i3\theta} \right) \left(1 - \frac{r^3}{R^3} e^{i3\theta} \right)^{-1} \right|$$

On effectue un développement limité à l'ordre r^3/R^3 soit :

$$\left| \frac{R^3 + \underline{Z}^3}{R^3 - \underline{Z}^3} \right| \simeq \left| 1 + 2 \frac{r^3}{R^3} e^{i3\theta} \right| = \sqrt{\left(1 + 2 \frac{r^3}{R^3} \cos 3\theta \right)^2 + 4 \frac{r^6}{R^6} \sin^2 \theta} \simeq 1 + 2 \frac{r^3}{R^3} \cos 3\theta$$

On reporte dans le logarithme et on termine le développement limité :

$$V(P) = \frac{\lambda}{2\pi\epsilon_0} \ln \left(1 + 2 \frac{r^3}{R^3} \cos 3\theta \right) \simeq \frac{\lambda}{\pi\epsilon_0} \frac{r^3}{R^3} \cos 3\theta$$

Cette expression vérifie l'expression déterminée avec les symétries en prenant $B_p = 0$ pour tout p et $A_p = 0$ si $p > 1$.

6. On ne peut pas utiliser l'expression de V précédente valable au voisinage de O . Il faut reprendre l'expression de départ.

Pour les électrodes impaires, les points de bord d'électrode dans la direction de O ont pour coordonnées $\underline{Z} = (R - a)$, $\underline{Z} = j(R - a)$ et $j^2 \underline{Z} = (R - a)$. On a donc :

$$R^3 - \underline{Z}^3 = R^3 - (R - a)^3 = R^3 - R^3 \left(1 - \frac{a}{R} \right)^3 \simeq R^3 - R^3 + 3aR^2 \simeq 3aR^2$$

De façon similaire, on obtient :

$$R^3 + \underline{Z}^3 = R^3 + (R - a)^3 = R^3 + R^3 \left(1 - \frac{a}{R} \right)^3 \simeq 2R^3 - 3aR^2$$

Le potentiel est donc :

$$V_0 = \frac{\lambda}{2\pi\epsilon_0} \ln \left(\frac{3a}{2R - 3a} \right) \simeq \frac{\lambda}{2\pi\epsilon_0} \ln \frac{3a}{2R}$$

Pour les électrodes paires, les valeurs de $\underline{Z}$ sont opposées aux précédentes, cela revient à inverser la fraction dans le logarithme et donc le potentiel des électrodes paires est $-V_0$.

7. La différence de potentiel entre les deux électrodes du condensateur ainsi formé est $\Delta V = -2V_0$ (car V_0 est négatif) et la charge par unité de longueur est $Q = 3\lambda$ puisqu'il y a trois fils par armatures. Ainsi la capacité $C = Q/\Delta V$ est $C = -3\lambda/2V_0$ ou $C = -3\pi\epsilon_0/\ln \left(\frac{3a}{2R} \right)$. On trouve alors que le potentiel au centre peut s'écrire :

$$V(r) = \frac{-2CV_0}{3\pi\epsilon_0} \frac{r^3}{R^3} \cos 3\theta$$

8. L'application numérique donne $C = 4,4 \cdot 10^{-11}$ F.

Chapitre 46

- A.1** 1. On introduit une autre origine O' avec la relation de Chasles :

$$\vec{p} = \sum q_i (\vec{OO'} + \vec{O'S_i}) = (\sum q_i) \vec{OO'} + \sum q_i \vec{O'S_i}$$

or $\sum q_i = 0$ donc la définition est indépendante de l'origine.

2. Dans le cas de deux charges :

$$\vec{p} = -q \vec{OS_1} + q \vec{OS_2} = q \vec{S_1S_2}$$

3. a) D'après l'hypothèse, l'hydrogène a perdu son électron donc il ne lui reste qu'un proton (charge positive du dipôle). Quant au fluor il a 9 protons et 10 électrons, c'est la charge négative du dipôle.
 b) La charge q du dipôle est la charge élémentaire e et le moment dipolaire, dirigé de F vers H, vaut $p = ed$ soit $p = 1,47 \cdot 10^{-29}$ C.m ou 4,42 D.
 c) La charge positive est à la position H et la charge négative en G. On prend F comme origine, alors :

$$p = eHF - 10eGF \Rightarrow GF = \frac{eHF - p}{10e} = 5,39 \cdot 10^{-12} \text{ m}$$

4. a) On applique le théorème de superposition :

$$V = \frac{q}{4\pi\epsilon_0} \left(\frac{-1}{r_1} + \frac{1}{r_2} \right)$$

- b) On effectue un développement limité au premier ordre :

$$\vec{S_1M} = \vec{S_1O} + \vec{OM} \Rightarrow S_1M = \sqrt{\left(\frac{a}{2}\right)^2 + r^2 + 2\frac{ar}{2} \cos \theta} \Rightarrow \frac{1}{r_1} \simeq \frac{1}{r} \left(1 - \frac{a}{2r} \cos \theta \right)$$

et :

$$\vec{S_2M} = \vec{S_2O} + \vec{OM} \Rightarrow S_2M = \sqrt{\left(\frac{a}{2}\right)^2 + r^2 - 2\frac{ar}{2} \cos \theta} \Rightarrow \frac{1}{r_2} \simeq \frac{1}{r} \left(1 + \frac{a}{2r} \cos \theta \right)$$

donc :

$$V_d(M) = \frac{qa \cos \theta}{4\pi\epsilon_0 r^2} = \frac{\vec{p} \cdot \vec{r}}{4\pi\epsilon_0 r^3}$$

5. a) On utilise les coordonnées sphériques :

$$\text{grad} \frac{1}{r^3} = \frac{d(r^{-3})}{dr} \vec{u}_r = \frac{-3}{r^4} \vec{u}_r = \frac{-3}{r^5} \vec{r}$$

et

$$\text{grad} \vec{p} \cdot \vec{r} = \text{grad} (pr \cos \theta) = p \cos \theta \frac{dr}{dr} \vec{u}_r + \frac{pr}{r} \frac{d \cos \theta}{d \theta} \vec{u}_\theta = p \cos \theta \vec{u}_r - p \sin \theta \vec{u}_\theta$$

ce qui est égal à $\vec{p}$.

- b) Pour calculer $\vec{E}$, on utilise $\vec{E} = -\text{grad} V_d$:

$$\vec{E} = \frac{1}{4\pi\epsilon_0} \text{grad} \frac{\vec{p} \cdot \vec{r}}{r^3} = \frac{1}{4\pi\epsilon_0} \left((\vec{p} \cdot \vec{r}) \text{grad} \frac{1}{r^3} + \frac{1}{r^3} \text{grad} (\vec{p} \cdot \vec{r}) \right)$$

En remplaçant par les expressions établies à la question précédente, on trouve le résultat demandé avec $k_f = 3$.

- c) En coordonnées sphériques, on a $\vec{r} = r \vec{u}_r$ et $\vec{p} = p(\cos \theta \vec{u}_r - \sin \theta \vec{u}_\theta)$. L'expression du champ devient :

$$E = \frac{1}{4\pi\epsilon_0 r^5} (3pr^2 \cos \theta \vec{u}_r - pr^2 \cos \theta \vec{u}_r + r^2 \sin \theta \vec{u}_\theta)$$

soit :

$$E_r = \frac{p \cos \theta}{2\pi\epsilon_0 r^3}; \quad E_\theta = \frac{p \sin \theta}{4\pi\epsilon_0 r^3}; \quad E_\phi = 0$$

d) On peut écrire $\tan \beta = E_\theta / E_r$, ce qui donne $\tan \beta = \frac{\tan \theta}{2}$.

e) Lorsque E est parallèle à Oy , on a $\theta + \beta = \frac{\pi}{2}$ donc :

$$\tan \beta = \tan \left(\frac{\pi}{2} - \theta \right) = \frac{1}{\tan \theta}$$

d'où $\tan^2 \theta = 2$ et $\theta = 54,7^\circ$.

6. a) Une surface équipotentielle est une surface sur laquelle tous les points sont au même potentiel. L'équipotentielle V_0 est telle que :

$$V_0 = \frac{p}{4\pi\epsilon_0} \frac{\cos \theta}{r^2} \Rightarrow r^2 = r_0^2 \cos \theta \text{ avec } r_0^2 = \frac{p}{4\pi\epsilon_0 V_0}$$

- b) Une ligne de champ est une ligne en tout point tangente au champ électrique. Si on appelle $d\vec{\ell}$ le vecteur déplacement élémentaire sur la ligne de champ, qui s'exprime par $d\vec{\ell} = dr\vec{u}_r + r d\theta\vec{u}_\theta$, l'équation s'obtient par :

$$\vec{E} \wedge d\vec{\ell} = \vec{0} \Rightarrow \frac{dr}{E_r} = \frac{r d\theta}{E_\theta} \Rightarrow \frac{dr}{r} = \frac{2 \cos \theta d\theta}{\sin \theta}$$

On trouve après intégration : $r = R \sin^2 \theta$, où R est une constante caractéristique de chaque ligne.

- c) Les équipotentielles et les lignes de champ sont tracées dans le cours.
7. a) Le couple exercé sur le dipôle par rapport à O est :

$$\vec{\Gamma} = \vec{OS}_1 \wedge (-q\vec{E}_e) + \vec{OS}_2 \wedge (q\vec{E}_e) = q\vec{S}_1 \wedge \vec{S}_2 \wedge \vec{E}_e = \vec{p} \wedge \vec{E}_e$$

Il est indépendant du point où il est calculé.

- b) Si l'on note α l'angle entre $\vec{p}$ et $\vec{E}_e$, l'énergie potentielle vaut $U = -pE_e \cos \alpha$. Elle est donc minimale lorsque $\alpha = 0$, c'est la position d'équilibre stable avec le dipôle aligné sur le champ, dans le même sens. Pour $\alpha = \pi$, l'énergie est maximale, c'est la position d'équilibre instable avec le dipôle anticolinéaire à $\vec{E}$.

A.2 On a représenté la distribution sur la figure (S46.1).

En appliquant le théorème de superposition, le potentiel créé en M est :

$$V(M) = \frac{1}{4\pi\epsilon_0} \left(\frac{2q}{OM} - \frac{q}{PM} - \frac{q}{P'M} \right)$$

En faisant l'hypothèse que $a \ll r$ on va calculer le potentiel à grande distance en effectuant un développement limité au deuxième ordre en a/r car il est nul au premier ordre. On rappelle que

$$(1+x)^{-1/2} \simeq \left(1 - \frac{x}{2} + \frac{3}{8}x^2 \right)$$

Il vient :

$$\vec{PM} = \vec{PO} + \vec{OM} \Rightarrow PM = \sqrt{a^2 + r^2 - 2ar \cos \theta} = r \sqrt{1 - 2\frac{a}{r} \cos \theta + \frac{a^2}{r^2}}$$

On obtient donc pour le développement limité :

$$\frac{1}{PM} \simeq \frac{1}{r} \left(1 + \frac{a}{r} \cos \theta - \frac{a^2}{2r^2} + \frac{3}{8} \frac{4a^2}{r^2} \cos^2 \theta \right)$$

On procède de même pour $P'M$ ce qui donne :

$$P'M = r \sqrt{1 + 2\frac{a}{r} \cos \theta + \frac{a^2}{r^2}} \Rightarrow \frac{1}{P'M} \simeq \frac{1}{r} \left(1 - \frac{a}{r} \cos \theta - \frac{a^2}{2r^2} + \frac{3}{8} \frac{4a^2}{r^2} \cos^2 \theta \right)$$

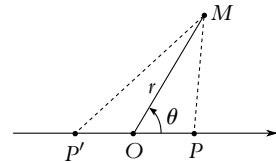


Figure S46.1

En ajoutant les trois termes qui forment le potentiel, on arrive à :

$$V = \frac{q}{4\pi\epsilon_0} \frac{a^2}{r^3} (1 - 3 \cos^2 \theta)$$

C'est ce qu'on appelle un développement quadripolaire car la charge totale est nulle ainsi que le moment dipolaire.

A.3 1. Pour exprimer l'énergie d'interaction entre ces deux molécules, il suffit d'exprimer l'énergie potentielle de l'une d'elles dans le champ électrostatique créé par l'autre : $E_p = -\vec{p}_1 \cdot \vec{E}_{2 \rightarrow 1}$ avec $\vec{E}_{2 \rightarrow 1} = \frac{1}{4\pi\epsilon_0} \frac{3(\vec{p}_2 \cdot \vec{P}_2 \vec{P}_1) \vec{P}_2 \vec{P}_1 - P_1 P_2^2 \vec{p}_2}{P_1 P_2^5}$.

D'où :

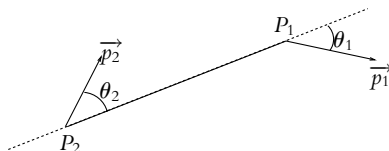
$$\begin{aligned} E_p &= -\frac{1}{4\pi\epsilon_0} \frac{3(\vec{p}_2 \cdot \vec{P}_2 \vec{P}_1)(\vec{p}_1 \cdot \vec{P}_2 \vec{P}_1) - P_1 P_2^2 \vec{p}_2 \cdot \vec{p}_1}{P_1 P_2^5} \\ &= \frac{1}{4\pi\epsilon_0} \frac{(P_1 P_2)^2 \vec{p}_2 \cdot \vec{p}_1 - 3(\vec{p}_2 \cdot \vec{P}_2 \vec{P}_1)(\vec{p}_1 \cdot \vec{P}_2 \vec{P}_1)}{(P_1 P_2)^5} \end{aligned}$$

On remarquera que cette expression est parfaitement invariante par échange des indices 1 et 2, ce qui est cohérent avec la notion d'interaction. On pourrait obtenir cette expression à partir de $E_p = -\vec{p}_2 \cdot \vec{E}_{1 \rightarrow 2}$.

2. Avec les notations proposées, l'énergie potentielle d'interaction s'écrit :

$$\begin{aligned} E_p &= \frac{1}{4\pi\epsilon_0} \frac{p_1 p_2 \cos(\theta_1 - \theta_2) - 3p_1 p_2 \cos \theta_1 \cos \theta_2}{(P_1 P_2)^3} \\ &= \frac{p_1 p_2}{4\pi\epsilon_0 (P_1 P_2)^3} (\sin \theta_1 \sin \theta_2 - 2 \cos \theta_1 \cos \theta_2) \end{aligned}$$

3.



On suppose que θ_1 fixe (on remarque qu'on aurait les mêmes résultats en prenant θ_2 fixe et en échangeant les indices).

On aura un équilibre lorsque $\frac{dE_p}{d\theta_2} = 0$.

$$\frac{dE_p}{d\theta_2} = \frac{p_1 p_2}{4\pi\epsilon_0 P_1 P_2^3} (\sin \theta_1 \cos \theta_2 + 2 \cos \theta_1 \sin \theta_2)$$

soit $\tan \theta_1 = -2 \tan \theta_2$.

4. L'équilibre sera stable si $\frac{d^2 E_p}{d\theta_2^2} > 0$. $\frac{d^2 E_p}{d\theta_2^2} = \frac{p_1 p_2}{4\pi\epsilon_0 P_1 P_2^3} (-\sin \theta_1 \sin \theta_2 + 2 \cos \theta_1 \cos \theta_2)$ d'où la condition $2 \cos \theta_1 \cos \theta_2 > \sin \theta_1 \sin \theta_2$.

5. On peut à partir de ces relations obtenir les positions d'équilibre particulières demandées :

a) $\theta_1 = 0$: l'équilibre est obtenu pour $\tan \theta_2 = -\frac{\tan \theta_1}{2} = 0$ soit

• $\theta_2 = 0$ alors $2 \cos \theta_1 \cos \theta_2 = 2 > \sin \theta_1 \sin \theta_2 = 0$ donc l'équilibre est stable,

—————> —————> stable

• $\theta_2 = \pi$ alors $2 \cos \theta_1 \cos \theta_2 = -2 < \sin \theta_1 \sin \theta_2 = 0$ donc l'équilibre est instable.

—————> —————< instable

b) $\theta_1 = \frac{\pi}{2}$: l'équilibre est obtenu pour $\tan \theta_2 = -\frac{\tan \theta_1}{2} = -\infty$ soit

- $\theta_2 = -\frac{\pi}{2}$ alors $2 \cos \theta_1 \cos \theta_2 = 0 > \sin \theta_1 \sin \theta_2 = -1$ donc l'équilibre est stable,



stable

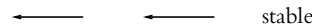
- $\theta_2 = -\frac{\pi}{2}$ alors $2 \cos \theta_1 \cos \theta_2 = 0 < \sin \theta_1 \sin \theta_2 = 1$ donc l'équilibre est instable.



instable

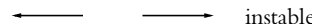
c) $\theta_1 = \pi$: l'équilibre est obtenu pour $\tan \theta_2 = -\frac{\tan \theta_1}{2} = 0$ soit

- $\theta_2 = \pi$ alors $2 \cos \theta_1 \cos \theta_2 = 2 > \sin \theta_1 \sin \theta_2 = 0$ donc l'équilibre est stable,



stable

- $\theta_2 = 0$ alors $2 \cos \theta_1 \cos \theta_2 = -2 < \sin \theta_1 \sin \theta_2 = 0$ donc l'équilibre est instable.



instable

d) $\theta_1 = -\frac{\pi}{2}$: l'équilibre est obtenu pour $\tan \theta_2 = -\frac{\tan \theta_1}{2} = \infty$ soit

- $\theta_2 = \frac{\pi}{2}$ alors $2 \cos \theta_1 \cos \theta_2 = 0 > \sin \theta_1 \sin \theta_2 = -1$ donc l'équilibre est stable,



stable

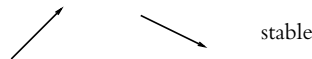
- $\theta_2 = \frac{\pi}{2}$ alors $2 \cos \theta_1 \cos \theta_2 = 0 < \sin \theta_1 \sin \theta_2 = 1$ donc l'équilibre est instable.



instable

e) $\theta_1 = \frac{\pi}{4}$: l'équilibre est obtenu pour $\tan \theta_2 = -\frac{\tan \theta_1}{2} = -\frac{1}{2}$ soit

- $\theta_2 = \text{Arctan}(-\frac{1}{2}) \simeq -26,6^\circ$ alors $2 \cos \theta_1 \cos \theta_2 \simeq 1,265 > \sin \theta_1 \sin \theta_2 \simeq -0,316$ donc l'équilibre est stable,



stable

- $\theta_2 = 180^\circ + \text{Arctan}(-\frac{1}{2}) \simeq 153,4^\circ$ alors $2 \cos \theta_1 \cos \theta_2 \simeq -1,265 < \sin \theta_1 \sin \theta_2 \simeq 0,316$ donc l'équilibre est instable.



instable

A.4 1. Le point A est placé sur l'axe du dipôle D_0 ce qui correspond à $\theta = 0$ avec l'utilisation habituelle des coordonnées sphériques. Le champ créé par D_0 en un point de cet axe est :

$$\vec{E} = \frac{P_0}{2\pi\epsilon_0 r^3} \vec{u}_r$$

On modélise le dipôle D_1 par une charge q à $r + a/2$ de O et une charge $-q$ à $r - a/2$ de O . La force s'exerçant sur la charge q est $\vec{F}_q = \frac{qP_0}{2\pi\epsilon_0(r+a/2)^3}\vec{u}_r$ et sur la charge $-q$, $\vec{F}_{-q} = \frac{-qP_0}{2\pi\epsilon_0(r-a/2)^3}\vec{u}_r$

2. La force totale s'exerçant sur le dipôle D_1 est :

$$\vec{F}_q + \vec{F}_{-q} = \frac{qP_0}{2\pi\epsilon_0 r^3} \left(\left(1 + \frac{a}{2r}\right)^{-3} - \left(1 - \frac{a}{2r}\right)^{-3} \right) \vec{u}_r \simeq -\frac{3qaP_0}{2\pi\epsilon_0 r^4} \vec{u}_r$$

Le moment dipolaire du dipôle induit est $P_1 = qa = \alpha\epsilon_0 E$. On en déduit :

$$\vec{F} = -\frac{3P_0^2\alpha}{4\pi^2\epsilon_0 r^7} \vec{u}_r$$

- B.1** 1. La réponse est 2. On a invariance par translation le long de Oz et par rotation autour de Oz donc on utilise les coordonnées cylindriques et on a indépendance vis-à-vis de z et ρ . Les plans contenant Oz et perpendiculaires à Oz sont des plans de symétrie et on n'a pas de plans d'antisymétrie donc le champ $\vec{E}$ est suivant $\vec{u}_\rho$. On applique le théorème de Gauss à un cylindre centré sur Oz de rayon ρ et de hauteur h et on obtient $\vec{E} = \frac{\lambda}{2\pi\epsilon_0\rho} \vec{u}_\rho$.

2. La réponse est 1. Le potentiel s'obtient par intégration soit

$$V = - \int \vec{E}(\rho) \cdot d\rho \vec{u}_\rho = -\frac{\lambda}{2\pi\epsilon_0} \ln \rho + \text{constante}$$

3. La réponse est 3. En appliquant le théorème de superposition, on obtient $V(M) = \frac{\lambda}{2\pi\epsilon_0} \ln \frac{O_1M}{O_2M}$ (la constante est nulle). En effectuant un développement limité, on a $V(M) = \frac{p \cos \theta}{2\pi\epsilon_0 r}$.
4. La réponse est 4. On a une équipotentielle si $r = C \cos \theta$ avec C une constante. Les équipotentielles sont des cercles centrés sur Ox dans le plan perpendiculaire à cet axe.
5. La réponse est 3. On obtient le champ par dérivation soit

$$\vec{E} = -\vec{\text{grad}} V = \frac{p \cos \theta}{2\pi\epsilon_0 r^2} \vec{u}_r + \frac{p \sin \theta}{2\pi\epsilon_0 r^2} \vec{u}_\theta$$

Son module est $\frac{p}{2\pi\epsilon_0 r^2}$. Il est constant si r est constant donc on a des cercles centrés en O et

$$\tan(\alpha - \theta) = \frac{E_\theta}{E_r} = \tan \theta \text{ soit } \alpha = 2\theta.$$

- B.2** 1. L'énergie d'interaction entre G_i et G_k est par exemple $U = -\vec{p}_k \cdot \vec{E}_i$ qui est égale à $-\vec{p}_i \cdot \vec{E}_k$. On a représenté les trois cas d'orientation possibles des dipôles sur la figure (S46.2). Sur les schémas de gauche et du centre l'angle θ est tel que $2\theta + \beta = \pi$ soit $\theta = 35,27^\circ$. C'est α qui a cette valeur sur le schéma de droite. On note r la distance entre les différents centres des dipôles égale à $2d \sin(\beta/2)$.

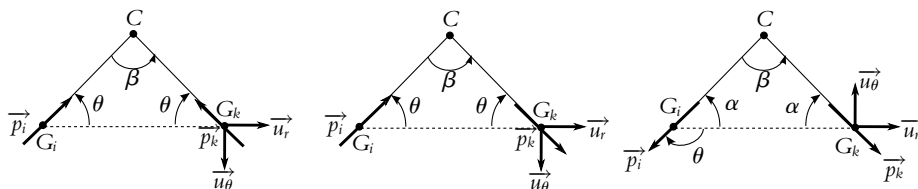


Figure S46.2

Dans le premier cas (cas de gauche), le champ créé par G_i à la position de G_k est : $\vec{E}_i = \frac{pk}{r^3}(2 \cos \theta \vec{u}_r + \sin \theta \vec{u}_\theta)$ et $\vec{p}_k = -p(\cos \theta \vec{u}_r + \sin \theta \vec{u}_\theta)$. L'énergie d'interaction est :

$$U_1 = -\frac{p^2 k}{r^3}(2 \cos^2 \theta + \sin^2 \theta)$$

Dans le deuxième cas (cas du centre), le champ créé par G_i à la position de G_k est encore : $\vec{E}_i = \frac{pk}{r^3}(2 \cos \theta \vec{u}_r + \sin \theta \vec{u}_\theta)$ et $\vec{p}_k = p(\cos \theta \vec{u}_r + \sin \theta \vec{u}_\theta)$. L'énergie d'interaction est :

$$U_2 = \frac{p^2 k}{r^3}(2 \cos^2 \theta + \sin^2 \theta)$$

Dans le troisième cas (cas de droite), le dipôle i n'est plus dans le même sens. L'expression générale du champ est la même mais avec $\theta = \pi - \alpha$ (figure S46.2). Le champ créé par G_i à la position de G_k est : $\vec{E}_i = \frac{pk}{r^3}(-2 \cos \alpha \vec{u}_r + \sin \alpha \vec{u}_\theta)$ et $\vec{p}_k = p(\cos \alpha \vec{u}_r - \sin \alpha \vec{u}_\theta)$. L'énergie d'interaction est :

$$U_3 = \frac{p^2 k}{r^3}(-2 \cos^2 \alpha - \sin^2 \alpha) = U_1$$

2. Il faut faire attention à ne compter qu'une fois l'énergie d'interaction entre deux dipôles. On peut procéder de la manière suivante : compter les trois interactions de G_1 avec les trois autres, puis celles de G_2 avec G_3 et G_4 , puis enfin celle de G_3 avec G_4 . Il y a cinq possibilités d'orientation :

- Les quatre dipôles vers le centre, alors $U = 6U_1$.
- Trois vers le centre et un opposé, alors $U = 3U_1 + 3U_2$.
- Deux vers le centre et deux opposés, alors $U = U_1 + 4U_2 + U_3$.
- Un vers le centre et trois opposés, alors $U = 3U_2 + 3U_3$.
- Quatre opposés, alors $U = 6U_3$.

3. Le point C est sur l'axe des dipôles. Le potentiel en C est $\pm \frac{kp}{d^2}$ suivant l'orientation du dipôle et donc l'énergie d'interaction avec la charge est $\pm \frac{kpe}{d^2}$ pour chaque dipôle. Si l'on note U' la valeur absolue de l'énergie d'interaction d'un dipôle avec la charge, pour les cinq cas précédents, l'énergie d'interaction totale est donc :

- Les quatre dipôles vers le centre, alors $E = 6U_1 + 4U'$.
- Trois vers le centre et un opposé, alors $E = 3U_1 + 3U_2 + 2U'$.
- Deux vers le centre et deux opposés, alors $E = U_1 + 4U_2 + U_3$.
- Un vers le centre et trois opposés, alors $E = 3U_2 + 3U_3 - 2U'$.
- Quatre opposés, alors $U = 6U_3 - 4U'$.

4. L'application numérique pour les cinq cas dans le même ordre donne :

- $E = 2,96^{-19}$ J.
- $E = 1,98^{-19}$ J.
- $E = 3,33^{-19}$ J.
- $E = -1,98^{-19}$ J.
- $E = -4,96^{-19}$ J.

5. L'arrangement le plus stable est celui de plus basse énergie donc le numéro (5).

6. Si on multiplie l'énergie trouvée pour le cas (5) par le nombre d'Avogadro, cela donne 299 kJ.mol^{-1} soit 25% d'erreur, ce qui n'est pas si mal pour une modélisation assez simple.

B.3 1. Le moment dipolaire global est nul avant le déplacement. Après le déplacement, la variation du moment dipolaire pour une paire d'ions est :

$$d\vec{p} = (e\delta_+ - e\delta_-)\vec{u}_x = ex\vec{u}_x$$

donc globalement $\vec{P} = Nex\vec{u}_x$.

2. a) À la position d'équilibre, un ion sodium est soumis à deux forces : la force électrique et une autre force f qui l'empêche de s'éloigner d'avantage (force de rappel).
 b) La relation d'équilibre projetée sur Ox entraîne :

$$0 = eE + f \Rightarrow f = -eE = -\frac{eP}{\epsilon_0 \chi_{\text{ion}}} = -\frac{Ne^2}{\epsilon_0 \chi_{\text{ion}}} x$$

3. a) La projection de la relation fondamentale de la dynamique en absence de champ pour les deux ions donne :

$$\begin{cases} m_+ \ddot{\delta}_+ = -Kx = -K(\delta_+ - \delta_-) \\ m_- \ddot{\delta}_- = Kx = -K(\delta_- - \delta_+) \end{cases}$$

- b) La somme des deux équations entraîne $m_+ \ddot{\delta}_+ + m_- \ddot{\delta}_- = 0$, ce qui prouve que le mouvement du centre de gravité de l'ensemble est une translation uniforme (ou l'immobilité). Pour avoir le mouvement relatif des ions par rapport au centre de gravité, on pourrait passer dans le référentiel barycentrique et utiliser le mobile fictif. On peut aussi combiner les deux équations en divisant la première par m_+ et en soustrayant la seconde divisée par m_- , ce qui donne :

$$\ddot{x} = -\frac{K}{m} x \Rightarrow \ddot{x} + \omega_T^2 x = 0$$

où $\omega_T = \sqrt{K/m}$ avec m la masse réduite des deux ions. Soit finalement :

$$\omega_T = \sqrt{\frac{Ne^2}{m\epsilon_0 \chi_{\text{ion}}}}$$

Chapitre 48

- A.1** 1. Tout plan passant par M et contenant le fil est un plan de symétrie : le champ magnétique est perpendiculaire à ce plan. En coordonnées cylindriques d'axe le fil, le champ magnétique est donc orthoradial.

On pourrait également dire que tout plan passant par M et perpendiculaire au fil est plan d'antisymétrie : le champ magnétique appartient à ce plan. L'information est cependant moindre que la précédente, on retiendra qu'il est préférable de chercher les plans de symétrie quand on veut déterminer la direction d'un champ magnétique.

2. La réponse et les arguments sont les mêmes qu'à la question précédente.
 3. Tout plan passant par M et perpendiculaire au fil est un plan d'antisymétrie : le champ magnétique appartient à ce plan. En coordonnées cylindriques d'axe le fil, le champ magnétique n'aura donc pas de composante axiale.

Pour un point M contenu dans le plan défini par les deux fils, on peut ajouter que ce plan est plan de symétrie des sources. Le champ magnétique est donc perpendiculaire à ce plan. En coordonnées cylindriques d'axe parallèle aux fils, le champ magnétique est donc orthoradial.

L'étude des cas particuliers donne :

- a) Si $I_1 = I_2$ alors le plan médiateur des deux fils est un plan de symétrie.
 b) Si $I_1 = -I_2$ alors le plan médiateur des deux fils est un plan d'antisymétrie.
 4. Tout plan passant par M et contenant l'axe Δ est un plan d'antisymétrie : le champ magnétique appartient à ce plan. En coordonnées cylindriques d'axe Δ , le champ magnétique n'aura donc pas de composante orthoradiale.
 Le plan perpendiculaire à Δ passant par le centre de la sphère est un plan de symétrie : le champ magnétique en un point M de ce plan est donc perpendiculaire à ce plan. En coordonnées cylindriques d'axe parallèle aux fils, le champ magnétique sera dirigé selon l'axe Δ .
 5. Tout plan perpendiculaire à la nappe et parallèle à l'axe (Oy) est un plan de symétrie : le champ magnétique en tout point d'un tel plan est donc perpendiculaire à ce dernier. Il est dirigé selon l'axe (Oz) .

6. Tout plan perpendiculaire à la nappe et perpendiculaire à l'axe (Oy) est plan d'antisymétrie : le champ magnétique en tout point de ce plan est contenu dans ce dernier. Il n'a donc pas de composante suivant (Oy).

Le plan (xOy) (en supposant que la largeur de la nappe s'étend entre $-\frac{l}{2}$ et $+\frac{l}{2}$) est plan de symétrie : le champ magnétique est perpendiculaire en tout point de ce plan à ce dernier. Il est dirigé selon (Oz) uniquement pour les points du plan (xOy). Pour les autres points, on sait juste que le champ magnétique n'a pas de composante suivant (Oy).

- A.2** 1. On a établi qu'un fil parcouru par un courant d'intensité crée un champ magnétique $\vec{B}_1 = \frac{\mu_0 I_1}{2\pi r} \vec{u}_\theta$. Une portion du deuxième fil subit donc la force de Laplace : $d\vec{F} = I_2 d\vec{l}_2 \wedge \vec{B}_1$ dirigé du fil 2 vers le fil 1.
On peut faire le même raisonnement en inversant le rôle des deux fils.
On a donc finalement une interaction attractive.

2. Le raisonnement est le même qu'à la question précédente en inversant le signe de I_2 : l'interaction est répulsive.
3. D'un point de vue magnétique, l'interaction est de même nature qu'à la question 1 à savoir attractive. Il faut cependant tenir compte de l'interaction électrique qui est répulsive. On ne peut pas conclure sur la nature de la résultante des interactions à partir des seules informations fournies car il faudrait pouvoir comparer les intensités des deux forces magnétique et électrique.
4. D'un point de vue magnétique, l'interaction est de même nature qu'à la question 2 à savoir répulsive. Il faut comme à la question précédente tenir compte de l'interaction électrique qui est répulsive. On a donc une interaction globale répulsive.
5. On retrouve une situation analogue à celle de la question 2 car le sens du courant est opposé au mouvement des électrons.
6. On retrouve une situation analogue à celle de la question 1 car le sens du courant est opposé au mouvement des électrons.

B.1 Soit un point P de la demi-spire et un élément de courant $d\vec{\ell}_P = a d\theta \vec{u}_\theta$ autour de P . La loi de Biot et Savart s'écrit :

$$\vec{B}(M) = \frac{\mu_0 I}{4\pi} \int_{\theta=-\pi/2}^{\theta=\pi/2} \frac{d\vec{\ell}_P \wedge \vec{PM}}{PM^3} = \frac{\mu_0 I}{4\pi} \int_{\theta=-\pi/2}^{\theta=\pi/2} \frac{d\vec{\ell}_P \wedge (\vec{PO} + \vec{OM})}{PM^3}$$

On cherche la composante du champ suivant Oz , or $d\vec{\ell} \wedge \vec{OM}$ est perpendiculaire à Oz alors que $d\vec{\ell} \wedge \vec{OP} = a d\ell \vec{u}_z$. L'intégration sur la demi-spire donne :

$$\vec{B}_z(M) = \frac{\mu_0 I}{4\pi} \frac{\pi a^2}{PM^3} \vec{u}_z = \frac{\mu_0 I}{4} \frac{a^2}{(a^2 + z^2)^{3/2}} \vec{u}_z$$

- B.2** 1. On applique la loi de Biot et Savart : $\vec{B}(O) = \frac{\mu_0}{4\pi} \int_{P \in \text{hélice}} Id\vec{l}_P \wedge \frac{\vec{PO}}{PO^3}$.
Or en coordonnées cartésiennes : $d\vec{l}_P = -R \sin \alpha d\alpha \vec{u}_x + R \cos \alpha d\alpha \vec{u}_y + \frac{h}{2\pi} d\alpha \vec{u}_z$,
 $\vec{PO} = -R \cos \alpha \vec{u}_x - R \sin \alpha \vec{u}_y - \frac{h\alpha}{2\pi} \vec{u}_z$ et $PO = \sqrt{R^2 + \left(\frac{h\alpha}{2\pi}\right)^2}$.

On en déduit :

$$\vec{B} = \frac{\mu_0 I}{4\pi} \int_1^2 \frac{1}{\left(R^2 + \left(\frac{h\alpha}{2\pi}\right)^2\right)^{\frac{3}{2}}} \left(\frac{Rh d\alpha}{2\pi} (\sin \alpha - \alpha \cos \alpha) \vec{u}_x - \frac{Rh d\alpha}{2\pi} (\alpha \sin \alpha + \cos \alpha) \vec{u}_y + R^2 d\alpha \vec{u}_z \right)$$

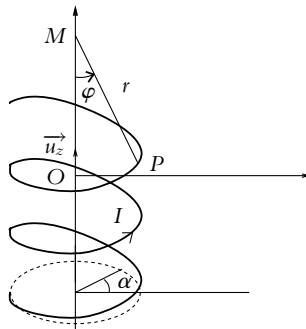
On ne calcule que la composante sur Oz soit :

$$B_z = \frac{\mu_0 I}{4\pi} \int_1^2 \frac{R^2}{\left(R^2 + \left(\frac{h\alpha}{2\pi}\right)^2\right)^{\frac{3}{2}}} d\alpha.$$

On utilise le changement de variables proposées et on obtient :

$d\alpha = -\frac{2\pi R}{h} \frac{d\varphi}{\sin^2 \varphi}$ et $R^2 + \left(\frac{h\alpha}{2\pi}\right)^2 = \frac{R^2}{\sin^2 \varphi}$. L'intégrale cherchée est donc :

$$\begin{aligned} B_z &= -\frac{\mu_0 I}{4\pi} \int_1^2 \frac{R^2}{\left(\frac{R^2}{\sin^2 \varphi}\right)^{\frac{3}{2}}} \frac{2\pi R}{h} \frac{d\varphi}{\sin^2 \varphi} = -\frac{\mu_0 I}{2h} \int_1^2 \sin \varphi d\varphi \\ &= \frac{\mu_0 I}{2h} (\cos \varphi_2 - \cos \varphi_1) \end{aligned}$$



- Dans le cas d'un solénoïde infini, $\alpha_1 \rightarrow -\infty$ soit $\varphi_1 \rightarrow -\pi$ et $\alpha_2 \rightarrow +\infty$ soit $\varphi_2 \rightarrow 0$. On en déduit : $B_z = \frac{\mu_0 I}{h}$.
- h est le pas de l'hélice ou la distance entre deux spires successives le long de l'axe quand on a tourné de 2π . On en déduit que $\frac{1}{h}$ est n le nombre de spires par unité de longueur du solénoïde et $B_z = \mu_0 n I$.
- Le champ magnétique possède également des composantes suivant (Ox) et (Oy) auxquelles on ne s'est pas intéressé ici. Le problème ne présente pas de symétrie permettant de conclure que le champ magnétique sur l'axe soit dirigé le long de cet axe. D'autre part, on ne s'est intéressé qu'au champ magnétique en un point de l'axe. On ne tirera donc pas de conclusion hâtive de ce résultat.
- Quand h tend vers 0, φ_1 tend vers φ_0^+ et φ_1 tend vers φ_0^- . On en déduit :

$$\cos \varphi_2 = \cos \varphi_0 - \sin \varphi_0 (\varphi_2 - \varphi_0) \quad \text{et} \quad \cos \varphi_1 = \cos \varphi_0 - \sin \varphi_0 (\varphi_1 - \varphi_0),$$

d'où : $\cos \varphi_2 - \cos \varphi_1 = \sin \varphi_0 (\varphi_1 - \varphi_2)$. On en déduit : $B_z = \frac{\mu_0 I \sin \varphi_0 (\varphi_1 - \varphi_2)}{2h}$.

$$\text{D'autre part, } \alpha_1 = \frac{2\pi R}{h} \cotan \varphi_1 = \frac{2\pi R}{h} \left(\cotan \varphi_0 - \frac{1}{\sin^2 \varphi_0} (\varphi_1 - \varphi_0) \right)$$

et de même $\alpha_2 = \frac{2\pi R}{h} \left(\cotan \varphi_0 - \frac{1}{\sin^2 \varphi_0} (\varphi_2 - \varphi_0) \right)$ d'où $\alpha_2 - \alpha_1 = \frac{2\pi R}{h \sin^2 \varphi_0} (\varphi_1 - \varphi_2)$ soit

$$\frac{\varphi_2 - \varphi_1}{h} = \sin^2 \varphi_0 \frac{\alpha_2 - \alpha_1}{2\pi R}. \text{ Or le nombre de spires est : } N = \frac{\alpha_2 - \alpha_1}{2\pi}. \text{ On retrouve donc le champ créé par } N \text{ spires sur leur axe } B(\varphi_0) = \frac{\mu_0 N I}{2R} \sin^2 \varphi_0.$$

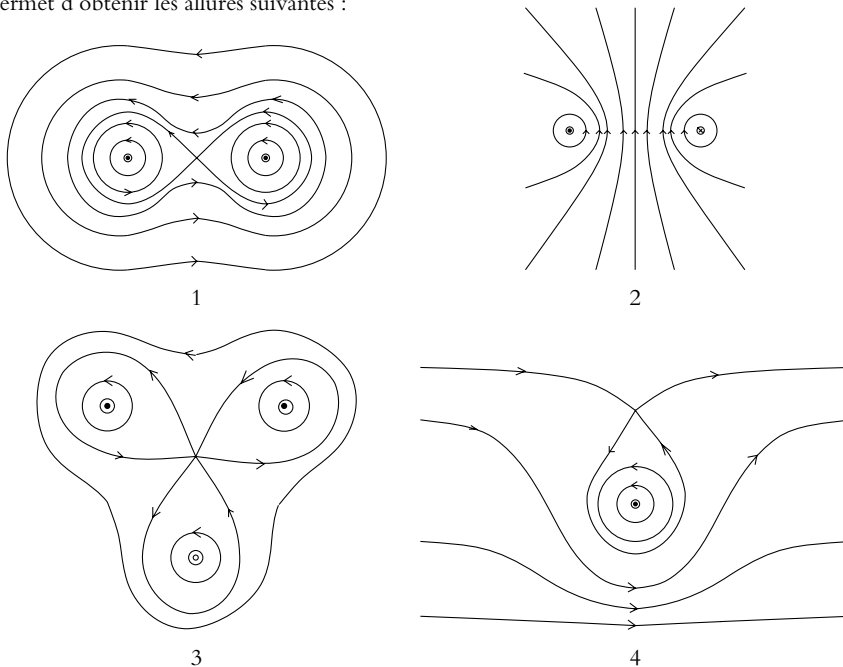
Chapitre 49

A.1 Une ligne de champ est tangente en tout point au champ magnétique. Lorsque deux lignes de champ se croisent, cela correspond à au moins deux orientations possibles pour le champ en un même point. La seule solution envisageable est la nullité du champ en ce point.

Dans le cas d'un champ magnétique, ce sera par exemple le cas à mi-chemin entre deux fils parcourus par des courants de même sens et de même intensité. Le champ créé par chacun des fils a la même valeur absolue mais une orientation opposée.

A.2 On utilise la propriété que les lignes de champ tournent autour du fil qui les crée ainsi que le résultat de l'exercice précédent.

Cela permet d'obtenir les allures suivantes :

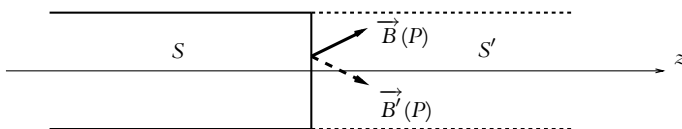


A.3 Les lignes de champ correspondent à celles d'un champ magnétique si le flux à travers toute surface fermée est nulle : le champ magnétique est à flux conservatif. C'est le cas des lignes de champ (a), (b), (c) et (e) qui correspondent respectivement au champ magnétique créé par une nappe de courants, au champ magnétique créé par un fil, à un champ magnétique uniforme et au champ magnétique créé par quatre fils.

Pour les cas (d) et (f), il suffit de calculer le flux à travers une sphère dont le centre est celui des lignes de champ pour se convaincre que le flux n'est pas nul donc ces lignes de champ ne sont pas celles d'un champ magnétique.

B.1 1. On appelle S' le solénoïde qui complète S en un solénoïde infini. S' se déduit de S par symétrie par rapport au plan de la face terminale de S .

Le champ $\vec{B}(P)$ créé par S en P et le champ $\vec{B}'(P)$ créé par S' en P sont donc symétriques par rapport à l'axe (Oz).



Par superposition, $\vec{B}(P) + \vec{B}'(P) = \mu_0 n I \vec{u}_z$. On en déduit, en projetant cette relation sur Oz :

$$B_z(P) = \frac{1}{2} \mu_0 n I.$$

2. a) À l'intérieur du solénoïde, celui-ci apparaît comme infini, les lignes de champ sont donc parallèles à l'axe Oz .
- b) On considère la surface fermée constituée des cercles de rayon x et y (perpendiculaires à Oz) et du tube de champ qui s'appuie sur ces deux cercles. Le flux de $\vec{B}$ à travers cette surface est nul d'où $-\mu_0 n I \pi x^2 + 0 + \frac{1}{2} \mu_0 n I \pi y^2 = 0$ donc $y = \sqrt{2}x$.

B.2 On montrera au chapitre suivant que le champ magnétique créé par une spire circulaire le long de son axe noté (Oz) vaut : $\vec{B} = \frac{\mu_0 I \sin^3 \alpha}{2R} \vec{u}_z$ avec $\cotan \alpha = \frac{z}{R}$.

On veut donc calculer : $\mathcal{C} = \int_{-\infty}^{+\infty} \frac{\mu_0 I \sin^3 \alpha}{2R} dz$. Or $dz = d(R \cotan \alpha) = -\frac{R d\alpha}{\sin^2 \alpha}$ et $\lim_{z \rightarrow +\infty} \alpha = 0$,

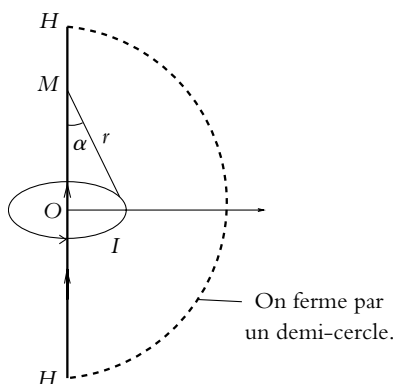
$\lim_{z \rightarrow -\infty} \alpha = \pi$. On en déduit : $\mathcal{C} = - \int_{\pi}^0 \frac{\mu_0 I \sin \alpha}{2} d\alpha = \mu_0 I$.

Ce résultat fait penser au théorème d'Ampère. Cependant l'axe en lui-même ne constitue pas un contour fermé : on le « ferme à l'infini » de la manière décrite sur la figure ci-dessous. On considère le demi-cercle représenté en pointillés sur la figure ci-contre et on fait tendre H vers l'infini. Par le théorème d'Ampère, la circulation le long de ce contour vaut :

$$\mathcal{C} = \mu_0 I_{\text{enlacé}} = \lim_{H \rightarrow \infty} \int_{-H}^H B dz + \int_{\text{demi-cercle}} \vec{B} \cdot d\vec{l}$$

avec $I_{\text{enlacé}} = I$.

Sur le demi-cercle, quand H tend vers l'infini, la spire se comporte comme un dipôle (Cf. chapitre ultérieur correspondant) et le champ y est donc en $\frac{1}{H^3}$ pour une longueur de πH soit finalement une contribution qui tend vers 0 pour H tendant vers l'infini. Le calcul pouvait donc être évité.



B.3 1. a) Tous les plans qui contiennent l'axe Oz sont plans d'antisymétrie de la distribution de courant. Le champ magnétique appartenant à tous ces plans, il est porté par $\vec{u}_z$. Il prend donc la forme : $\vec{B}(M) = B(z) \vec{u}_z$.

b) Le plan de la spire est plan de symétrie de la distribution de courant. $\vec{B}$ étant un pseudo-vecteur, $\vec{B}(-z) = \vec{B}(z)$ donc $B(-z) = B(z)$.

c) Soit un point P de la spire et un élément de courant $d\vec{\ell}_P = R d\theta \vec{u}_\theta$ autour de P . La loi de Biot et Savart s'écrit :

$$\vec{B}(M) = \frac{\mu_0 I}{4\pi} \int_{\theta=-\pi}^{\theta=\pi} \frac{d\vec{\ell}_P \wedge \vec{PM}}{PM^3} = \frac{\mu_0 I}{4\pi} \int_{\theta=-\pi}^{\theta=\pi} \frac{d\vec{\ell}_P \wedge (\vec{PO} + \vec{OM})}{PM^3}$$

On cherche la composante du champ suivant Oz , or $d\vec{\ell} \wedge \vec{OM}$ est perpendiculaire à Oz alors que $d\vec{\ell} \wedge \vec{OP} = R d\vec{\ell} \vec{u}_z$. L'intégration sur la spire donne :

$$B(z) = \frac{\mu_0 I}{4\pi} \frac{2\pi R^2}{PM^3} = \frac{\mu_0 I}{2} \frac{R^2}{(R^2 + z^2)^{3/2}}$$

On peut le mettre sous la forme $B(z) = B_0 f(u)$ avec $B_0 = B(0) = \frac{\mu_0 I}{2R}$ et $f(u) = (1 + u^2)^{-3/2}$ avec $u = z/R$.

2. a) Le plan contenant P et l'axe (Oz) est plan d'antisymétrie de la distribution de courant donc $\vec{B}(P)$ appartient à ce plan, il n'a pas de composante orthoradiale.

b) La distribution de courant est invariante par toute rotation d'axe (Oz) donc $\vec{B}(P) = \vec{B}(r, z)$.

c) Le flux de $\vec{B}$ à travers une surface fermée est nul.

d) Au voisinage de l'axe, il n'y a aucun courant. Le théorème d'Ampère appliqué à un contour voisin de l'axe montre que la circulation de $\vec{B}$ le long de ce contour est nulle.

- e) Le rayon r étant petit, on considère B_z uniforme sur les faces du cylindre pour le calcul du flux de $\vec{B}$. De plus, B_r est uniforme sur la surface latérale de ce cylindre. Le flux magnétique à travers une surface fermée est nul donc :

$$-B_z(0, z)\pi r^2 + B_r(r, z)2\pi r dz + B_z(0, z + dz)\pi r^2 = 0$$

Un développement limité au premier ordre en r permet d'obtenir la relation demandée :

$$B_r(r, z) = -\frac{r}{2} \frac{dB_z}{dz}(0, z)$$

- f) Le contour envisagé n'entourant aucun courant, la circulation de $\vec{B}$ le long de ce contour est nulle. On a donc :

$$-B_r(r, z)dr - B_z(r + dr, z)dz + B_r(r, z + dz)dr + B_z(r, z)dz = 0$$

Un développement limité au premier ordre en $drdz$ donne : $\frac{\partial B_r}{\partial z} = \frac{\partial B_z}{\partial r}$. En utilisant le résultat de la question précédente et en intégrant la relation obtenue entre $r = 0$ et r , on en déduit :

$$B_z(r, z) = B(0, z) - \frac{r^2}{4} \frac{d^2 B}{dz^2}(0, z).$$

- g) En utilisant les résultats précédents et les données de l'énoncé, on peut exprimer le champ magnétique dans le plan de la spire ($z = 0$) : $B_z(r, 0) = B_0 \left(1 + \frac{3r^2}{4R^2}\right)$. On cherche r tel que $\frac{B_z(r, 0) - B_0}{B_0} \leq 10^{-2}$, on obtient : $r \leq \frac{0,2}{\sqrt{3}}R = 0,115R$.

Chapitre 50

- A.1** 1. Le système est invariant par translation selon les axes (Ox) et (Oy) : le champ magnétique n'est donc fonction que de z .

Il n'y a pas d'invariance par rotation.

2. Tout plan parallèle au plan (xOz) est plan de symétrie des sources. En tout point, le champ magnétique est donc perpendiculaire à ces plans : il est dirigé selon (Oy) soit $\vec{B}(M) = B_y(z)\vec{u}_y$.
Le plan (xOy) est également un plan de symétrie ; les champs magnétiques en un point d'altitude z et en un point d'altitude $-z$ vérifient : $\vec{B}(z) = -\vec{B}(-z)$ soit $B_y(z) = -B_y(-z)$.

3. et 4. On se place en $z > 0$. On applique le théorème d'Ampère sur le contour orienté représenté sur la figure ci-contre :

$$\int_{\text{contour}} \vec{B} \cdot d\vec{l} = B_y(z)l - B_y(-z)l = 2B_y(z)l$$

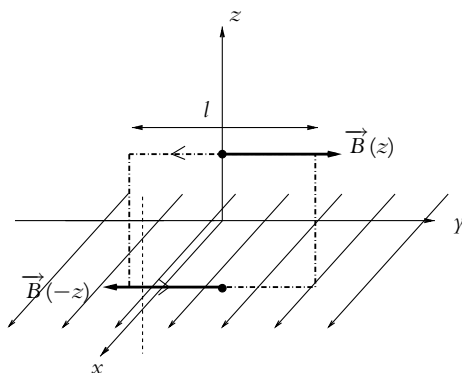
et d'après le théorème d'Ampère :

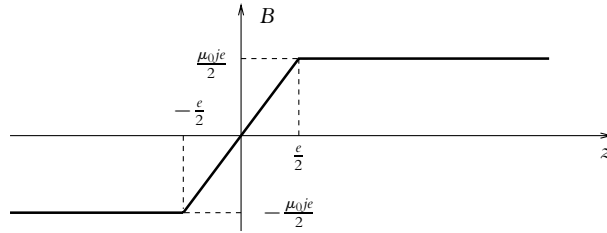
$$\int_{\text{contour}} \vec{B} \cdot d\vec{l} = \mu_0 I_{\text{enlacé}} \quad \text{dont le calcul donne :}$$

$$\mu_0 I_{\text{enlacé}} = \begin{cases} \mu_0 j 2z l & \text{si } 0 < z < \frac{e}{2} \\ \mu_0 j e l & \text{si } z \leq \frac{e}{2} \end{cases}$$

On en déduit :

$$B_y(z) = \begin{cases} \mu_0 j z & \text{si } |z| < \frac{e}{2} \\ \text{signe}(z) \frac{\mu_0 j e}{2} & \text{si } |z| \leq \frac{e}{2} \end{cases}$$





A.2 La distribution est invariante par rotation d'axe Oz et par translation selon Oz donc le champ $\vec{B}$ ne dépend que de r .

Tout plan perpendiculaire à l'axe Oz est plan de symétrie de la distribution de courant, donc $\vec{B}$ est selon Oz . En conclusion $\vec{B} = B(r)\vec{u}_z$.

On choisit comme contours, des rectangles dont deux des côtés sont selon Oz (figure S50.1).

Le vecteur circulation élémentaire en N_1 et N_2 est perpendiculaire à $\vec{u}_z$ donc à $\vec{B}$; la circulation sur ces côtés est nulle. Finalement, la circulation est égale à :

$$C = \int_A^B \vec{B}(M_1) \cdot d\vec{\ell}_{M_1} + \int_C^D \vec{B}(M_2) \cdot d\vec{\ell}_{M_2} = \ell(B(M_1) - B(M_2))$$

si l'on note $\ell = AB = CD$.

Pour calculer le courant enlacé, il y a plusieurs cas dessinés sur la figure (S50.2). On a grisé les zones où le courant enlacé n'est pas nul.

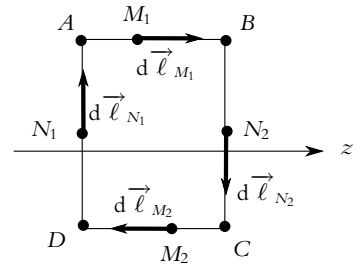


Figure S50.1

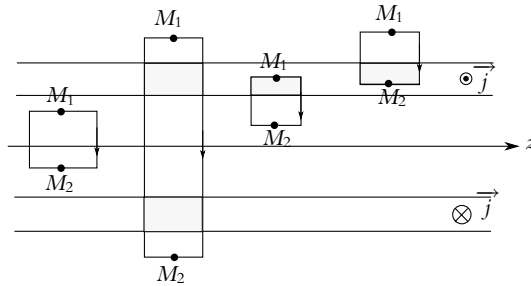


Figure S50.2

Premier cas, le contour est entièrement à l'intérieur du solénoïde : $I_{\text{enlacé}} = 0$, donc $B(M_1) = B(M_2)$. Le champ magnétique est uniforme à l'intérieur du solénoïde et on le note B_i .

Deuxième cas, les deux points sont à l'extérieur : $I_{\text{enlacé}} = -j(b-a)\ell + j(b-a)\ell = 0$ donc on a encore $B(M_1) = B(M_2)$. Le champ magnétique est uniforme à l'extérieur du solénoïde et on le note B_e .

Troisième cas, le point M_1 est à une distance r de l'axe, dans la zone de courant : $I_{\text{enlacé}} = -j(r-a)\ell$, le signe «-» est dû au fait que $\vec{j}$ est dans le sens négatif sur le dessin par rapport au sens d'orientation du contour. On a $B(M_1) = B_i + \mu_0 j(a-r)$.

Enfin quatrième cas, c'est M_2 qui est dans la zone de courant : $I_{\text{enlacé}} = -j(b-r)\ell$. On a $B_e = B(M_2) - \mu_0 j(b-r)$.

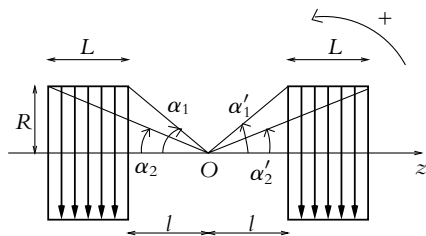
Il subsiste donc deux inconnues B_e et B_i . On utilise une propriété admise pour le solénoïde infini : le champ extérieur B_e est nul (une autre méthode serait de calculer le champ sur l'axe par empilement de spires). On trouve alors :

$$B_e = 0; \quad \vec{B}_i = \mu_0 j(b-a)\vec{u}_z; \quad \vec{B} = \mu_0 j(r-a)\vec{u}_z \quad \text{pour } a \geq r \geq b$$

A.3 1. On va utiliser le calcul du champ magnétique créé par un solénoïde effectué dans le cours :

$$B_z = \frac{\mu_0 n I}{2} (\cos \alpha_2 - \cos \alpha_1).$$

On a ici deux solénoïdes soit $B_z = \frac{\mu_0 n I}{2} (\cos \alpha_2 - \cos \alpha_1 + \cos \alpha'_2 - \cos \alpha'_1)$. Or $\cos \alpha_1 = \cos \alpha'_1 = \frac{l}{\sqrt{R^2 + l^2}}$ et $\cos \alpha_2 = \cos \alpha'_2 = \frac{l+L}{\sqrt{R^2 + (l+L)^2}}$.



On en déduit : $B_z = \mu_0 n I \left(\frac{l+L}{\sqrt{R^2 + (l+L)^2}} - \frac{l}{\sqrt{R^2 + l^2}} \right)$ soit

$$k = \mu_0 n \left(\frac{l+L}{\sqrt{R^2 + (l+L)^2}} - \frac{l}{\sqrt{R^2 + l^2}} \right)$$

Application numérique : on trouve $k = 2,76 \cdot 10^{-4}$. Pour une intensité de quelques ampères, on aura donc un champ de l'ordre du millitesla, ce qui est une valeur usuelle de champ magnétique.

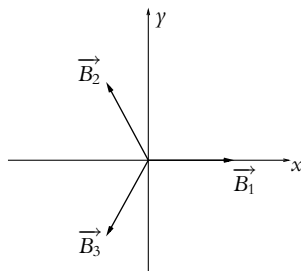
2. On a la superposition de trois systèmes analogues au précédent décalés l'un par rapport à l'autre d'un angle de $\frac{2\pi}{3}$.

Le champ magnétique résultant est :

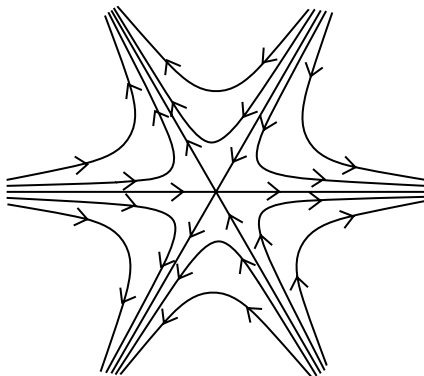
$$\vec{B} = \vec{B}_1 + \vec{B}_2 + \vec{B}_3 = \vec{0}$$

car sur (Ox) : $B_x = B - B \cos \frac{2\pi}{3} - B \cos \frac{2\pi}{3} = 0$ et

$$B_y = 0 + B \sin \frac{2\pi}{3} - B \sin \frac{2\pi}{3} = 0.$$



Les lignes de champ ont l'allure suivante :



3. On a maintenant :

$$\begin{cases} i_1 = I\sqrt{2} \cos \omega t \\ i_2 = I\sqrt{2} \cos \left(\omega t - \frac{2\pi}{3} \right) \\ i_3 = I\sqrt{2} \cos \left(\omega t - \frac{4\pi}{3} \right) \end{cases}$$

On en déduit :

$$\vec{B} = kI\sqrt{2} \begin{vmatrix} \cos \omega t + \cos \left(\omega t - \frac{2\pi}{3} \right) \cos \frac{2\pi}{3} + \cos \left(\omega t - \frac{4\pi}{3} \right) \cos \frac{4\pi}{3} \\ \cos \left(\omega t - \frac{2\pi}{3} \right) \sin \frac{2\pi}{3} + \cos \left(\omega t - \frac{4\pi}{3} \right) \sin \frac{4\pi}{3} \end{vmatrix}$$

En utilisant les relations trigonométriques usuelles, on trouve :

$$\vec{B} = \frac{3}{\sqrt{2}} kI (\cos \omega t \vec{u}_x + \sin \omega t \vec{u}_y)$$

On a donc un champ tournant à la pulsation ω et de module $\frac{3}{\sqrt{2}} kI$. Ce dispositif est utilisé dans les moteurs (Cf. cours de deuxième année PSI).

B.1 1. a) La distribution de courant est invariante par rotation d'axe Oz , donc $\vec{B}$ est indépendant de θ . La distribution est invariante par translation selon Oz donc $\vec{B}$ est indépendant de z . Le champ $\vec{B}$ ne dépend que de r .

b) Le plan contenant P et l'axe Oz est plan de symétrie de la distribution ; le champ $\vec{B}$ est donc perpendiculaire à ce plan et $\vec{B}$ est suivant $\vec{u}_\theta$.

2. a) Étant donnée la forme de $\vec{B}$, on choisit comme contour d'Ampère des cercles de rayon r et d'axe Oz . Le vecteur déplacement élémentaire sur un cercle est $d\vec{\ell} = d\ell \vec{u}_\theta$. La circulation de $\vec{B}$ sur ce contour fermé s'écrit :

$$C = \oint_C \vec{B}_P \cdot d\vec{\ell}_P = \oint_C B(r) \vec{u}_\theta \cdot d\ell \vec{u}_\theta = \oint_C B(r) d\ell = 2\pi r B(r)$$

Pour un point P extérieur au cylindre, le courant enlacé est celui qui parcourt le fil, donc c'est le flux de $\vec{J}$ à travers le disque de rayon R . D'autre part $\vec{J}$ est dans le sens positif par rapport au contour orienté suivant $\vec{u}_\theta$, donc $I_{\text{enlacé}} = \pi R^2 J$. Le théorème d'Ampère entraîne :

$$2\pi r B = \mu_0 \pi R^2 J \Rightarrow \vec{B} = \mu_0 \frac{JR^2}{2r} \vec{u}_\theta$$

b) Dans le cas où $r < R$ seule la fraction du courant qui passe à travers le contour est à prendre en compte donc $I_{\text{enlacé}} = \pi r^2 J$. Le théorème d'Ampère entraîne :

$$2\pi r B = \mu_0 \pi r^2 J \Rightarrow \vec{B} = \mu_0 \frac{Jr}{2} \vec{u}_\theta$$

c) Avec $\vec{r} = r \vec{u}_r$ et $\vec{J} = J \vec{u}_z$, on calcule $\vec{J} \wedge \vec{r}$ ce qui donne $Jr \vec{u}_\theta$. le résultat précédent peut donc s'écrire :

$$\vec{B} = \frac{\mu_0}{2} \vec{J} \wedge \vec{r}$$

3. Les symétries et invariances sont les mêmes que dans les questions précédentes. Le circulation sur un cercle fermé est donc $C = 2\pi r B$.

a) Pour la zone $r < b_2$, aucun courant ne traverse le contour, donc $I_{\text{enlacé}} = 0$ et $\vec{B} = \vec{0}$.

- b) Pour la zone intermédiaire, on a grisé sur la figure (S50.3), la surface concernée pour le calcul de $I_{\text{enlacé}}$.

Le flux de $\vec{J}$ est $I_{\text{enlacé}} = J\pi(r^2 - b_2^2)$. Le théorème d'Ampère entraîne :

$$2\pi rB = \mu_0 J\pi(r^2 - b_2^2) \Rightarrow \vec{B} = \frac{\mu_0 J}{2} \left(r - \frac{b_2^2}{r} \right) \vec{u}_\theta$$

4. On va superposer le champ magnétique $\vec{B}_1$ créé par un cylindre d'axe O_1z , de rayon b_1 et de densité de courant $\vec{J}$ avec celui $\vec{B}_2$ créé par un cylindre d'axe O_2z , de rayon b_2 et de densité de courant $-\vec{J}$. La superposition des densités de courant donne bien la distribution proposée.

Un point P intérieur à la cavité est intérieur aux deux cylindres précédents. On utilise le résultat intrinsèque déterminé à la question (1) :

$$\vec{B}_1 = \frac{\mu_0}{2} \vec{J} \wedge \overrightarrow{O_1P} \text{ et } \vec{B}_2 = \frac{\mu_0}{2} (-\vec{J}) \wedge \overrightarrow{O_2P} \Rightarrow \vec{B} = \frac{\mu_0}{2} \vec{J} \wedge (\overrightarrow{O_1P} - \overrightarrow{O_2P})$$

On obtient un champ uniforme :

$$\vec{B} = \frac{\mu_0}{2} \vec{J} \wedge \overrightarrow{O_1O_2}$$

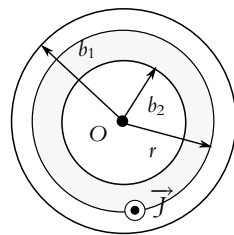


Figure S50.3

- B.2** 1. Les réponses (A) et (C) sont justes, les deux autres sont fausses.

2. Puisque le champ $\vec{B}$ est perpendiculaire aux plans de symétrie de la distribution de courant, il est orthoradial. Seule la réponse (A) est juste.

3. D'autre part, la distribution de courant est invariante par rotation d'axe Oy , donc le champ B ne dépend que de la distance à l'axe que l'on note r et de y . D'après ce qui précède, on choisit comme contour pour calculer la circulation des cercles de rayons r et d'axe Oy orienté dans le sens trigonométrique (positif par rapport à l'orientation de Oy). La circulation sur ces contours, auxquels $\vec{B}$ est tangent et le long desquels il a une norme constante est $C = 2\pi rB$.

Dans le cas d'un point M du plan xOy d'abscisse x , on a $r = x$. Le courant enlacé par un cercle intérieur au tore est celui des N spires (coté vertical à la distance $R - a$ de Oy). Le courant circule selon Oy donc dans le sens positif par rapport à l'orientation du contour. On a donc $I_{\text{enlacé}} = nI$ et le théorème d'Ampère s'écrit :

$$2\pi rB = \mu_0 nI \Rightarrow B = \frac{\mu_0 nI}{2\pi r} = \frac{\mu_0 nI}{2\pi x}$$

La réponse (B) est correcte.

4. Le flux du champ magnétique à travers la surface d'une spire est le même pour toutes les spires que pour celui du plan xOy , il s'écrit :

$$\Phi = \int \int \vec{B} \cdot \vec{dS} = \frac{\mu_0 nI}{2\pi} \int_{y=-a/2}^{y=a/2} \int_{x=R-a/2}^{x=R+a/2} \frac{dx dy}{x} = \frac{\mu_0 nIa}{2\pi} \ln \frac{2R+a}{2R-a}$$

La réponse (A) est juste.

5. On a :

$$B_{\text{max}} = \frac{\mu_0 nI}{2\pi \left(R - \frac{a}{2}\right)} \text{ et } B_{\text{min}} = \frac{\mu_0 nI}{2\pi \left(R + \frac{a}{2}\right)}$$

Le champ moyen B_{moy} est $(B_{\text{max}} + B_{\text{min}})/2$ et la variation relative recherchée est :

$$10\% = 0,1 = \frac{B_{\text{max}} - B_{\text{min}}}{B_{\text{moy}}} = 2 \frac{B_{\text{max}} - B_{\text{min}}}{B_{\text{max}} + B_{\text{min}}} \Rightarrow \frac{a}{R} = 0,100$$

- B.3** 1. Le champ créé par un fil infini a été calculé dans le cours. Si le fil est selon l'axe Oz , avec I dans le sens de l'axe, et que l'on repère la position d'un point M en coordonnées cylindriques (r, θ, z) ,

le champ magnétique en M est :

$$\vec{B}(M) = \frac{\mu_0 I}{2\pi r} \vec{u}_\theta$$

2. a) L'intensité dans tout le ruban est I_1 donc celle par unité de largeur est $I_1/(2a)$. Pour une largeur dy , l'intensité est $dI = \frac{I_1 dy}{2a}$.

- b) Le plan xOz est plan de symétrie de la distribution de courant, donc le champ en M est perpendiculaire à ce plan et selon Oy .

On a représenté le champ élémentaire créé par la largeur dy en M sur la figure (S50.4).

On commence par projeter ce champ sur la direction finale (attention B est dans le sens des y décroissants) soit :

$$dB_y = -\frac{\mu_0 dI}{2\pi r} \cos \alpha = -\frac{\mu_0 I_1 dy}{4\pi a} \frac{z}{(y^2 + z^2)}$$

L'intégration donne :

$$B_y = -\frac{\mu_0 I_1 z}{4\pi a} \int_{-a}^a \frac{dy}{(y^2 + z^2)} = -\frac{\mu_0 I_1}{4\pi a} \left[\arctan \frac{y}{z} \right]_{-a}^a = -\frac{\mu_0 I_1}{2\pi a} \arctan \frac{a}{z}$$

- c) Avec les sens opposés des courants, les deux champs s'ajoutent et donc le champ total est le double du champ précédent.

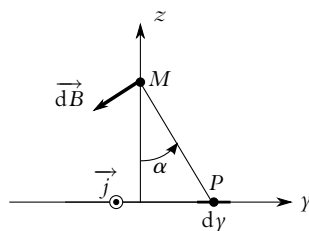


Figure S50.4

B.4 Disque de Rowland

1. La densité surfacique de courant au point P (figure S50.5) est $\vec{j}_s = \sigma \vec{v}_P$ soit $\vec{j}_s = \sigma r \omega \vec{u}_\theta$ si l'on utilise les coordonnées cylindriques d'axe Oz .

2. a) La spire élémentaire est circulaire, de rayon compris entre r et $r + dr$. Le courant dI dans cette spire est $j_s dr$.

- b) D'après le cours, le champ créé en un point M de l'axe (figure S50.6) par la spire élémentaire est :

$$dB_z = \frac{\mu_0 dI}{2r} \sin^3 \alpha$$

Or on a $\tan \alpha = r/z$ donc $dr = \frac{z d\alpha}{\cos^2 \alpha}$. Le champ s'écrit donc :

$$\begin{aligned} dB_z &= \frac{\mu_0 \sigma \omega z}{2} \frac{\sin^3 \alpha}{\cos^2 \alpha} d\alpha \\ &= \frac{\mu_0 \sigma \omega z}{2} \left(\frac{\sin \alpha}{\cos^2 \alpha} d\alpha - \cos \alpha d\alpha \right) \end{aligned}$$

L'intégration pour $\alpha \in [0, \alpha_m]$ avec $\alpha_m = \arctan(R/z)$ donne :

$$B_z = \frac{\mu_0 \sigma \omega z}{2} \left(\frac{1}{\cos \alpha_m} - 1 - \sin \alpha_m \right)$$

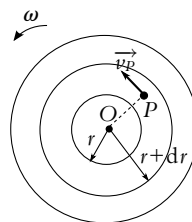


Figure S50.5

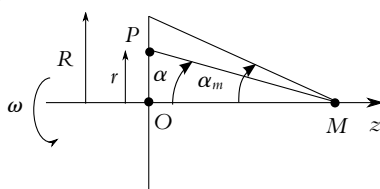


Figure S50.6

- B.5** 1. La longueur de bobinage ℓ est telle que $\sin \alpha = (r_2 - r_1)/\ell$ et le nombre de spires est $N = \ell/a$ donc la réponse (D) est juste.

2. En utilisant encore les relations dans le triangle, on peut écrire $dz = d\ell \cos \alpha$, or le nombre dN de spire a une longueur $d\ell = a dN$, d'où $dN = \frac{dz}{a \cos \alpha}$.

3. La résistance de $dN = \frac{d\ell}{a \sin \alpha}$ (soit $dN = \frac{dr}{a \sin \alpha}$) spires de rayon r (section $s = \pi a^2/4$) est :

$$dR = \frac{\rho 2\pi r}{s} dN = \frac{\rho 2\pi r dr}{sa \sin \alpha}$$

On intègre la relation entre r_1 et r_2 :

$$R = \frac{\rho 2\pi}{sa \sin \alpha} \int_{r_1}^{r_2} r dr = \frac{\rho \pi (r_2^2 - r_1^2)}{sa \sin \alpha}$$

En remplaçant s par son expressions on trouve la réponse (B).

4. Le calcul effectué dans le cours valide la réponse (A).
5. L'angle est le même pour toutes les spires. Pour un nombre dN de spires, le champ créé en S est :

$$d\vec{B} = \frac{\mu_0 I dN}{2r} \sin^3 \alpha \vec{u}_z = \frac{\mu_0 I dr}{2ar} \sin^2 \alpha \vec{u}_z$$

en intégrant entre r_1 et r_2 , on trouve :

$$\vec{B} = \frac{\mu_0 I}{2a} \sin^2 \alpha \ln \frac{r_2}{r_1} \vec{u}_z$$

La réponse (C) est juste.

- B.6** 1. La surface S s'appuyant sur le contour fermé orienté est elle-même orientée par rapport à l'orientation du contour, selon la règle de la main droite. La circulation du champ $\vec{B}$ sur le contour C est égal au flux du vecteur densité de courant $\vec{j}$ à travers la surface S multiplié par μ_0 .
2. a) Le plan contenant le point M et l'axe Oz est plan de symétrie de la distribution de courant donc le champ $\vec{B}(M)$ est orthogonal à ce plan : il est orthoradial.
b) La distribution de courant est invariante par translation selon Oz et par rotation d'axe Oz donc $\vec{B}$ ne dépend que de r .
c) Les lignes de champ sont donc des cercles d'axe Oz .
3. a) On calcule la circulation du champ le long d'un cercle d'axe Oz et de rayon r orienté dans le sens positif par rapport à Oz , c'est-à-dire suivant $\vec{u}_\theta$. Avec ce qui précède, il vient :

$$\oint_{\text{cercle}} \vec{B}(M) \cdot d\ell \vec{u}_\theta = 2\pi r B(r)$$

Si $r > R_3$, le courant enlacé par le contour est nul car égal à $I - I$. Le théorème d'Ampère donne alors :

$$2\pi r B(r) = \mu_0 I_{\text{enlacé}} = 0 \Rightarrow B = 0$$

- b) L'intérêt d'un câble coaxial par rapport à un câble normal est la nullité du champ à l'extérieur donc pas de parasites induits sur le reste du circuit.
4. L'intensité est le flux du vecteur densité de courant à travers la section du conducteur. Pour $r < R_1$, on a en norme $I = j_1 \pi R_1^2$ et pour $R_2 < r < R_3$, toujours en norme $I = j_2 \pi (R_3^2 - R_2^2)$ soit avec les orientations :

$$\vec{j}_1 = \frac{I}{\pi R_1^2} \vec{u}_z \text{ et } \vec{j}_2 = -\frac{I}{\pi (R_3^2 - R_2^2)} \vec{u}_z$$

5. Comme précédemment, on prend comme contour des cercles de rayon r et d'axe Oz , orientés dans le sens positif par rapport à Oz , c'est-à-dire suivant $\vec{u}_\theta$. La circulation de $\vec{B}$ sur ces contours est $2\pi r B$. On doit maintenant calculer le flux de $\vec{j}_1$ et $\vec{j}_2$ à travers la surface du contour :

- a) $r > R_1$, $I_{\text{enlacé}} = j_1 \pi r^2$ soit $B = \frac{\mu_0 j_1 r}{2}$.
- b) $R_1 < r < R_2$, $I_{\text{enlacé}} = j_1 \pi R_1^2$ soit $B = \frac{\mu_0 j_1 R_1^2}{2r}$.
- b) $R_2 < r < R_3$, $I_{\text{enlacé}} = j_1 \pi R_1^2 - j_2 \pi (r^2 - R_2^2)$ soit $B = \frac{\mu_0 j_1 R_1^2}{2r} - \frac{\mu_0 j_2 (r^2 - R_2^2)}{2r}$.

6. La vérification de la continuité est immédiate.
 7. La courbe $B(r)$ est tracée sur la figure (S50.7).

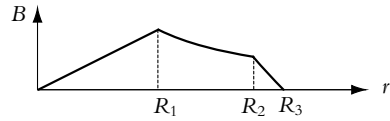
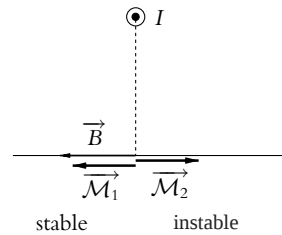


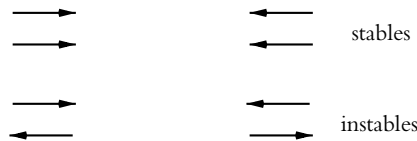
Figure S50.7

Chapitre 51

A.1 Le fil crée un champ magnétique $\vec{B} = \frac{\mu_0 I}{2\pi} \vec{u}_\theta$ en choisissant les coordonnées cylindriques d'axe le fil. L'énergie potentielle d'un aimant, qu'on modélise par un dipôle magnétique de moment $\vec{M}$, dans un champ magnétique extérieur s'exprime par : $E_p = -\vec{M} \cdot \vec{B}$. Cette énergie doit passer par un extremum pour que l'aimant soit en équilibre. Ce sera le cas lorsque le moment magnétique de l'aimant et le champ magnétique créé par le fil seront colinéaires. Comme le moment magnétique de l'aimant est horizontal, il faut que le champ magnétique le soit, ce qui n'est le cas que sur la verticale passant par le fil. Le champ magnétique est perpendiculaire au fil donc l'aimant doit également être perpendiculaire au fil. On a donc deux positions d'équilibre possibles : l'aimant doit être à la verticale du fil et perpendiculaire à la direction du fil. Les deux sens sont possibles : l'un conduit à une position d'équilibre stable quand champ et moment sont de même sens, l'autre à une position d'équilibre instable quand champ et moment sont de sens opposé.

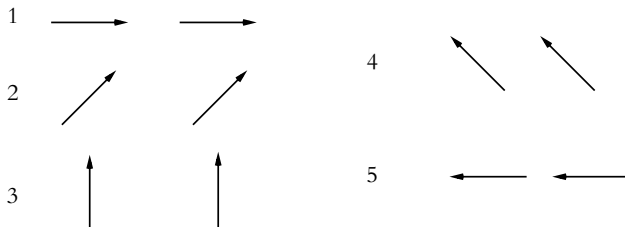


A.2 1. Les deux aimants créent chacun un champ magnétique dont les lignes de champ sont celles d'un dipôle magnétique. L'énergie potentielle du second aimant dans le champ créé par le premier s'écrit : $E_p = -\vec{M} \cdot \vec{B}$. Les positions d'équilibre correspondent aux extrema de ce produit scalaire : le moment magnétique d'un aimant et le champ magnétique créé par l'autre doivent être colinéaires. L'équilibre sera stable si le moment magnétique d'un aimant et le champ magnétique créé par l'autre sont de même sens. On en déduit les positions d'équilibre suivantes :



Toutes les orientations sont possibles, on n'a donné ici que les orientations relatives des deux aimants.

2. On applique le résultat de la question précédente aux situations proposées :



- B.1** 1. On a représenté le système sur la figure (S51.1). On utilise les coordonnées sphériques de centre O_1 . Le champ magnétique créé par le dipôle en un point M de la spire est :

$$B_r = \frac{\mu_0 M_1}{4\pi} \frac{2 \cos \theta}{r^3} \quad \text{et} \quad B_\theta = \frac{\mu_0 M_1}{4\pi} \frac{\sin \theta}{r^3} \quad \text{avec} \quad r = O_1 P.$$

Avec l'orientation prise pour i_2 sur le schéma (sens direct par rapport au repère), un élément de courant de la spire est $i_2 d\ell \vec{u}_\varphi$. La force de Laplace s'exerçant sur cet élément est donc :

$$d\vec{F} = i_2 d\ell \vec{u}_\varphi \wedge \vec{B} = i_2 d\ell (-B_\theta \vec{u}_r + B_r \vec{u}_\theta)$$

La symétrie du système impose une force suivant l'axe Oz , on se limite donc à calculer la composante projetée sur $\vec{u}_z$ de la force :

$$\begin{aligned} dF_z &= i_2 d\ell (-B_\theta \cos \theta - B_r \sin \theta) \\ &= i_2 d\ell \frac{\mu_0 M_1}{4\pi r^3} (-\sin \theta \cos \theta - 2 \sin \theta \cos \theta) \end{aligned}$$

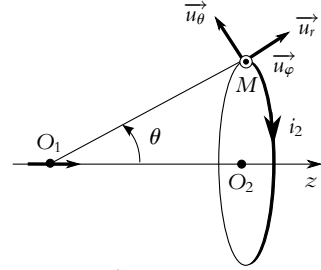


Figure S51.1

Tous les points de la spire sont à égale distance $r = \sqrt{R_2^2 + d^2}$ de O_1 . D'autre part, $d\ell = R_2 d\varphi$, donc finalement

$$F_z = -\frac{3\mu_0 M_1 R_2 i_2}{4\pi (R_2^2 + d^2)^{3/2}} \sin \theta \cos \theta \int_0^{2\pi} d\varphi = -\frac{3\mu_0 M_1 R_2 i_2}{2(R_2^2 + d^2)^{3/2}} \sin \theta \cos \theta$$

Si l'on remplace $\sin \theta$ et $\cos \theta$ par leurs expressions, on trouve :

$$F_z = -\frac{3\mu_0 i_2 M_1 R_2^2 d}{2(R_2^2 + d^2)^{5/2}}$$

qui est une force attractive. En supposant que $d \gg R_2$, on trouve :

$$F_z = -\frac{3\mu_0 i_2 M_1 R_2^2}{2d^4}$$

2. Le moment magnétique de S_2 est $\vec{M}_2 = \pi R_2^2 i_2 \vec{u}_z$. L'énergie potentielle de ce dipôle dans le champ de $\vec{M}_1$ est $E_p = -\vec{M}_2 \cdot \vec{B}_1$ où $\vec{B}_1$ est le champ créé par M_1 en O_2 ($\theta = 0$) soit :

$$E_p = -\frac{2\mu_0 M_1}{4\pi z^3} \pi R_2^2 i_2 = -\frac{\mu_0 M_1 i_2 R_2^2}{2z^3}$$

où z est la position de O_2 . Comme les dipôles ne sont pas induits, la force est $\vec{F} = -\vec{\text{grad}} E_p$ soit

$$F_z = -\frac{dE_p}{dz} = -\frac{3\mu_0 M_1 i_2 R_2^2}{2z^4}$$

Avec $z = d$ on retrouve bien le résultat précédent. Attention, il ne faut pas fixer la position à $z = d$ avant de dériver sinon on trouve une force nulle.

- B.2** 1. On se place en coordonnées polaires dans le plan méridien passant par M et on a $\theta = \frac{\pi}{2} - \lambda$. Par conséquent, le champ magnétique terrestre a pour composante dans ce plan :

$$\vec{B} = \frac{\mu_0 \mathcal{M}}{4\pi R^3} (2 \sin \lambda \vec{u}_r + \cos \lambda \vec{u}_\theta)$$

en notant R le rayon de la Terre.

2. $\tan I = \frac{B_r}{B_\theta} = 2 \tan \lambda.$

3. La composante horizontale du champ magnétique est la composante suivant $\vec{u}_\theta$ soit $B_h = \frac{\mu_0 \mathcal{M} \cos \lambda}{4\pi R^3}$ donc :

$$\mathcal{M} = \frac{4\pi R^3 B_h}{\mu_0 \cos \lambda}$$

L'application numérique fournit : $\mathcal{M} = 8,08.10^{22} \text{ A.m}^2$ et $I = 66,5^\circ$.

4. On obtient une valeur de I proche de celle observée : le modèle semble convenir. Il n'est cependant pas parfait puisqu'on a des écarts entre valeur mesurée et valeur théorique.

Chapitre 52

- A.1** 1. On a établi que le module de la vitesse d'une particule est constant lors de son mouvement dans un champ magnétique : l'énergie cinétique est donc conservée.
2. Si le module de la vitesse est constant, il n'en est pas de même de sa direction. Par conséquent, la quantité de mouvement peut changer de direction en conservant son module.
3. Le moment cinétique en O vaut : $\vec{L}_O = \vec{OM} \wedge m \vec{v} = mR^2 \omega \vec{u}_z$. Il est donc constant.

- A.2** 1. L'énergie cinétique s'écrit : $E_c = \frac{1}{2}mv^2$. Comme $m_e < m_p$, on en déduit que la vitesse de l'électron est plus grande que celle du proton : $v_e > v_p$.
2. Le rayon de la trajectoire s'exprime par : $R = \frac{mv}{|q|B}$ donc $\frac{R_e}{R_p} = \frac{m_e v_e}{m_p v_p} = \frac{E_{c_e} v_p}{E_{c_p} v_e} = \frac{v_p}{v_e}$. On en déduit que le rayon de la trajectoire de l'électron est plus petit que celui du proton : $R_e < R_p$.
3. La pulsation cyclotron vaut : $\Omega = \frac{|q|B}{m}$ donc la pulsation du mouvement de l'électron est plus grande que celle du mouvement du proton. Comme la pulsation est inversement proportionnelle à la période, la période de l'électron est plus petite que celle du proton : $T_e < T_p$.

- A.3** 1. a) L'accélération $\vec{a}$ est constante, soit :

$$\frac{d\vec{v}}{dt} = \vec{a} = cte \Rightarrow \vec{v} = \vec{a}t + \vec{v}_0$$

La norme de $\vec{v}$ varie donc l'énergie cinétique. Or seule la force électrique travaille (la force magnétique $q(\vec{v} \wedge \vec{B})$ est perpendiculaire à $\vec{v}$ donc à la trajectoire), le champ est un champ électrique.

Pour que la trajectoire soit rectiligne, il faut, d'après l'expression de $\vec{v}$ que $\vec{a}$ soit colinéaire à $\vec{v}_0$. Or $\vec{a} = \frac{q}{m}\vec{E}$, donc $\vec{E}$ est colinéaire à $\vec{v}_0$.

- b) On note O le centre du repère. On intègre l'expression de la vitesse pour avoir la position $\vec{OM}$ de la particule :

$$\vec{OM} = \frac{q}{2m}t^2\vec{E} + t\vec{v}_0 + \vec{OM}_0$$

M_0 étant la position initiale du point.

2. a) La trajectoire circulaire est la trajectoire d'une charge dans un champ magnétique perpendiculaire à la vitesse initiale. On en déduit que $\vec{B}$ est suivant Oz et que $\vec{v}_0$ est dans le plan xOy .
- b) On note $\omega = \dot{\theta}$ la vitesse angulaire. La trajectoire est circulaire, donc la vitesse en coordonnées polaires a pour coordonnées $\vec{v} = (0, R_0\omega, 0)$, et l'accélération se réduit à $\vec{a} = (-R_0\omega^2, R_0\dot{\omega})$, puisque $\dot{R}_0 = 0$. La relation fondamentale de la dynamique donne :

$$m\ddot{\vec{a}} = q(\vec{v} \wedge \vec{B})$$

soit en projection en polaires :

$$\vec{v} \wedge \vec{B} \begin{cases} R_0\omega B \\ 0 \end{cases} \Rightarrow \begin{cases} -mR_0\omega^2 = qR_0\omega B \\ R_0\dot{\omega} = 0 \end{cases}$$

Des deux équations on tire $\omega = -\frac{qB}{m} = cte$, donc si la charge est positive elle tourne dans le sens rétrograde (horaire) par rapport à Oz . Puisque ω est constante, le mouvement est circulaire uniforme et $v_0 = R_0|\omega|$ d'où, $R_0 = \frac{mv_0}{qB}$.

A.4 1. a) On écrit la relation fondamentale de la dynamique dans le référentiel galiléen :

$$m \vec{a} = q(\vec{E} + \vec{v} \wedge \vec{B})$$

Sachant que la vitesse à un instant quelconque a pour composante $(\dot{x}, \dot{y}, \dot{z})$, cela donne en projection sur les axes :

$$\vec{v} \wedge \vec{B} \begin{cases} 0 \\ \dot{z}B \\ -\dot{y}B \end{cases} \quad \text{soit} \quad \begin{cases} m\ddot{x} = qE \\ m\ddot{y} = q\dot{z}B \\ m\ddot{z} = -q\dot{y}B \end{cases}$$

- b) La première équation s'intègre en $m\dot{x} = qEt + cte_1$. La cte_1 est déterminée à $t = 0$, et puisque $\dot{x}(0) = 0$ elle est nulle. On intègre encore une fois pour obtenir $mx = \frac{qEt^2}{2} + cte_2$. À $t = 0$, $x = 0$, donc $cte_2 = 0$.
- c) Les deux autres équations sont des équations couplées, c'est-à-dire qui dépendent chacune des deux variables. Il y a deux méthodes pour résoudre ce genre de système : soit par substitution, ce que nous allons faire dans cette question, soit en utilisant les complexes, ce qui sera fait dans la question suivante.
- On intègre l'équation en $\ddot{y}$, soit : $m\dot{y} = qzB + cte_3$. À $t = 0$, $\dot{y} = 0$ et $z = 0$, donc $cte_3 = 0$. On reporte alors l'expression de $\dot{y}$ dans la dernière équation, ce qui donne :

$$\ddot{z} = -\left(\frac{qB}{m}\right)^2 z \Rightarrow \ddot{z} + \omega_c^2 z = 0$$

Il s'agit de l'équation d'un oscillateur harmonique, dont la solution est :

$$z = A_1 \cos \omega_c t + A_2 \sin \omega_c t$$

Pour déterminer les constantes A_1 et A_2 il faut calculer $\dot{z}$, soit :

$$\dot{z} = \omega_c(-A_1 \sin \omega_c t + A_2 \cos \omega_c t)$$

À l'instant $t = 0$, on obtient : $z(0) = 0 = A_1$ et $\dot{z}(0) = \omega_c A_2 = v_0$.

- d) Il reste à calculer y sachant que $\dot{y} = \omega_c z$, soit $\dot{y} = v_0 \sin \omega_c t$, on obtient en intégrant $y = -\frac{v_0}{\omega_c} \cos \omega_c t + cte_4$. À $t = 0$, $y = 0$ donc $cte_4 = \frac{v_0}{\omega_c}$. On récapitule les équations paramétriques de la trajectoire :

$$\begin{cases} x(t) = \frac{qE}{2m} t^2 \\ y(t) = \frac{v_0}{\omega_c} (1 - \cos \omega_c t) \\ z(t) = \frac{v_0}{\omega_c} \sin \omega_c t \end{cases}$$

- e) La projection de la trajectoire dans le plan yOz a pour équation :

$$\left(y - \frac{v_0}{\omega_c}\right)^2 + z^2 = \left(\frac{v_0}{\omega_c}\right)^2$$

Il s'agit donc de l'équation d'un cercle $\mathcal{C}$ de centre $(y_c = \frac{v_0}{\omega_c}, z_c = 0)$ et de rayon $R = \frac{v_0}{\omega_c}$. Quant au mouvement selon Ox il est uniformément accéléré. La composition des ces deux mouvements donne un hélice d'axe parallèle à Ox , de projection $\mathcal{C}$, et dont le pas p (variation de x en un tour) croît avec t : $p = \frac{qE}{2m} \left(\frac{2\pi}{\omega_c}\right)^2$.

2. a) On écrit la relation fondamentale de la dynamique dans le référentiel galiléen :

$$m \vec{a} = q(\vec{E} + \vec{v} \wedge \vec{B})$$

Sachant que la vitesse à un instant quelconque a pour composante $(\dot{x}, \dot{y}, \dot{z})$, cela donne en projection sur les axes :

$$\vec{v} \wedge \vec{B} \begin{cases} \dot{y}B \\ -\dot{x}B \\ 0 \end{cases} \quad \text{soit} \quad \begin{cases} m\ddot{x} = q\dot{y}B \\ m\ddot{y} = -q\dot{x}B + qE \\ m\ddot{z} = 0 \end{cases}$$

- b) On intègre facilement l'équation sur z qui donne $z = 0$ puisque $\dot{z}(0) = 0$ et $z(0) = 0$. Le mouvement a lieu dans le plan xOy .
- c) On utilise la méthode complexe ; pour cela on combine les deux équations restantes en multipliant la seconde par i , ce qui donne :

$$m(\ddot{x} + i\ddot{y}) = iqE + qB(\dot{y} - i\dot{x}) = iqE - i(\dot{x} + i\dot{y})$$

Sachant que $\underline{\dot{C}} = \dot{x} + i\dot{y}$ et $\underline{\ddot{C}} = \ddot{x} + i\ddot{y}$, l'équation devient :

$$\underline{\ddot{C}} + i\omega_c \underline{\dot{C}} = i \frac{qE}{m}$$

- d) Il s'agit d'une équation linéaire à coefficients constant du premier ordre en $\underline{\dot{C}}$. La solution générale de l'équation homogène est $\underline{\dot{C}}_H = \underline{A} \exp(-i\omega_c t)$ et sachant que le second membre est constant, une solution particulière (constante) est $\underline{\dot{C}}_p = \frac{qE}{\omega_c m} = \frac{E}{B}$.
- À $t = 0$, $\dot{x}(0) = v_0$ et $\dot{y}(0) = 0$, soit $\underline{\dot{C}} = v_0$, d'où :

$$v_0 = \underline{A} + \frac{E}{B} \Rightarrow \underline{\dot{C}} = \left(v_0 - \frac{E}{B}\right) \exp(-i\omega_c t) + \frac{E}{B}$$

- e) On intègre de nouveau pour avoir $\underline{C}$:

$$\underline{C} = -\frac{1}{i\omega_c} \left(v_0 - \frac{E}{B}\right) \exp(-i\omega_c t) + \frac{E}{B} t + cte$$

Sachant qu'à $t = 0$, $\underline{C} = 0$ on en déduit que $cte = \frac{1}{i\omega_c} \left(v_0 - \frac{E}{B}\right)$, d'où :

$$\underline{C} = \frac{i}{\omega_c} \left(v_0 - \frac{E}{B}\right) (e^{-i\omega_c t} - 1) + \frac{E}{B} t$$

Pour terminer on prend la partie réelle et la partie imaginaire :

$$\begin{cases} x(t) = \text{Re}(\underline{C}) = \frac{1}{\omega_c} \left(v_0 - \frac{E}{B}\right) \sin \omega_c t + \frac{E}{B} t \\ y(t) = \text{Im}(\underline{C}) = \frac{1}{\omega_c} \left(v_0 - \frac{E}{B}\right) (\cos \omega_c t - 1) \end{cases}$$

- f) On remarque que si $v_0 = \frac{E}{B}$, alors $x(t) = v_0 t$ et $y(t) = 0$. Donc seules les particules ayant cette vitesse initiale ont un mouvement rectiligne uniforme suivant l'axe Ox . Ce dispositif sert à sélectionner les particules en fonction de leur vitesse.

A.5 1. On néglige la force de pesanteur devant la force électrostatique. Si l'on note $\vec{E}$ le champ électrostatique, la force subie par la charge est $\vec{F} = q\vec{E}$. Dans le cas d'un champ uniforme, ce qui est le cas d'un condensateur plan, la relation générale $\vec{E} = -\vec{\text{grad}} V$ devient ici par intégration $\vec{E} = -\frac{U}{d} \vec{u}_x$. On en déduit :

$$\vec{F} = -q \frac{U}{d} \vec{u}_x$$

2. a) On applique la relation fondamentale de la dynamique à l'électron : $m\vec{a} = \vec{F}$. La force n'a pas de composante sur Oy et comme la vitesse initiale n'a pas de composantes non plus sur Oy on en déduit que le mouvement est dans le plan xOz . La projection de la relation fondamentale avec l'expression de $\vec{F} = \vec{E}$ et les intégrations successives donnent :

$$\begin{cases} m\ddot{x} = -q \frac{U}{d} \\ m\ddot{z} = 0 \end{cases} \Rightarrow \begin{cases} \dot{x} = -\frac{q}{m} \frac{U}{d} t (+cte_1) \\ \dot{z} = cte_2 = v_0 \end{cases} \Rightarrow \begin{cases} x = -\frac{q}{m} \frac{U}{d} \frac{t^2}{2} (+cte_3) \\ z = v_0 t (+cte_4) \end{cases}$$

Les constantes entre parenthèses (cte_1, cte_3, cte_4) sont déterminées nulles grâce aux conditions initiales.

En éliminant t entre les équations paramétriques en x et z , on obtient :

$$x = -\frac{q}{m} \frac{U}{d} \frac{z^2}{2v_0^2} = \frac{e}{m} \frac{U}{d} \frac{z^2}{2v_0^2}$$

Il s'agit d'une parabole.

- b) Le point de sortie correspond à $z_K = D$, soit dans l'équation précédente $x_K = \frac{e}{m} \frac{U}{d} \frac{D^2}{2v_0^2}$. D'après les équations paramétriques obtenues à la question (a), l'instant de passage en K est $t_K = z_K/v_0$ soit D/v_0 . les composantes de la vitesse en K sont obtenues en reportant t_K dans $\dot{x}$ et $\dot{z}$, soit :

$$\begin{cases} \dot{x}_K = \frac{e}{m} \frac{U}{d} \frac{D}{v_0} \\ \dot{z}_K = v_0 \end{cases}$$

- c) En dehors des plaques, puisqu'on néglige l'effet du champ de pesanteur et que l'on suppose le champ électrique nul, aucune force ne s'exerce sur l'électron : sa trajectoire est donc rectiligne uniforme (principe d'inertie dans un référentiel galiléen). On obtient aussi ce résultat par intégration de $\vec{A} = \vec{0}$.

- d) Dans le triangle O_eJP (figure S52.1), $X_P = O_eP = (JO_1 + L) \tan \theta$. Il faut donc exprimer $\tan \theta$ et JO_1 sachant que le segment JP est tangent à la parabole en K . On peut écrire :

$$\tan \theta = \frac{O_1K}{JO_1} = \frac{x_K}{JO_1} = \frac{e}{m} \frac{U}{d} \frac{D^2}{2v_0^2} \frac{1}{JO_1}$$

et

$$\tan \theta = \left(\frac{dx}{dz} \right)_K = \left(\frac{\dot{x}}{\dot{z}} \right)_K = \frac{e}{m} \frac{U}{d} \frac{D}{v_0^2}$$

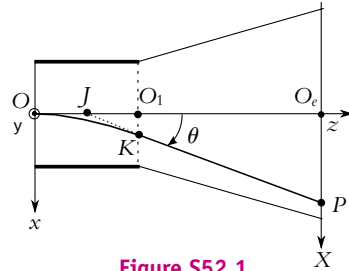


Figure S52.1

En combinant les deux équations, on trouve (c'est d'ailleurs une propriété de la parabole) que $JO_1 = D/2$ et donc :

$$X_P = \left(\frac{D}{2} + L \right) \frac{eD}{mdv_0^2} U$$

La déviation est proportionnelle à la tension U .

B.1 On applique la relation : $R = \frac{\int_A^B \vec{E} \cdot d\vec{l}}{\sigma \iint \vec{E} \cdot d\vec{S}}$ dans les configurations cylindrique et sphérique.

- La symétrie cylindrique implique que $\vec{E} = E(r)\vec{u}_r$ en coordonnées cylindriques. On en déduit : $\int_A^B \vec{E} \cdot d\vec{l} = V_1 - V_2 = \int_{R_1}^{R_2} E(r)dr$. Par ailleurs, l'intensité I est égale au flux de $\vec{j} = \sigma \vec{E}$ à la surface latérale du cylindre de rayon de r soit : $I = 2\pi rhj(r) = 2\pi r\sigma E(r)$. Alors la résistance dR entre deux cylindres de rayon r et $r + dr$ vérifie : $\vec{E} \cdot d\vec{l} = E(r)dr = -dV = dR I$ donc d'après les résultats précédents : $dR = -\frac{E(r)dr}{2\pi r\sigma E(r)} = -\frac{dr}{2\pi r\sigma}$ et en intégrant entre R_1 et R_2 : $R = \frac{1}{2\pi\sigma h} \ln \frac{R_2}{R_1}$.
- La symétrie sphérique implique que $\vec{E} = E(r)\vec{u}_r$ en coordonnées sphériques. On en déduit : $\int_A^B \vec{E} \cdot d\vec{l} = V_1 - V_2 = \int_{R_1}^{R_2} E(r)dr$. Par ailleurs, l'intensité I est égale au flux de $\vec{j} = \sigma \vec{E}$ à travers la sphère de rayon de r soit : $I = 4\pi r^2 j(r) = 4\pi r^2 \sigma E(r)$. Alors la résistance dR entre deux sphères de rayon r et $r + dr$ vérifie : $\vec{E} \cdot d\vec{l} = E(r)dr = -dV = dR I$ donc d'après les résultats précédents : $dR = -\frac{E(r)dr}{4\pi r^2 \sigma E(r)} = -\frac{dr}{4\pi r^2 \sigma}$ et en intégrant entre R_1 et R_2 : $R = \frac{R_2 - R_1}{4\pi\sigma R_1 R_2}$.

B.2 1. La force F_v est en N ou kg.m.s^{-2} . Le produit mv a pour unité des kg.m.s^{-1} , donc τ est homogène à un temps.

2. a) Les forces sont au nombre de trois : force électrique, force magnétique et force de frottement. On écrit la relation fondamentale de la dynamique dans le référentiel galiléen :

$$m\vec{a} = q(\vec{E} + \vec{v} \wedge \vec{B}) - \frac{m\vec{v}}{\tau}$$

En régime permanent, l'accélération est nulle, donc l'équation à laquelle obéit la vitesse est :

$$q(\vec{E} + \vec{v} \wedge \vec{B}) - \frac{m\vec{v}}{\tau} = 0$$

b) On calcule le produit vectoriel et on projette l'équation précédente sur les trois axes :

$$\vec{v} \wedge \vec{B} \begin{cases} \dot{y}B \\ -\dot{x}B \\ 0 \end{cases} \quad \text{soit} \quad \begin{cases} 0 = q\dot{y}B - \frac{m}{\tau}\dot{x} + qE_x \\ 0 = -q\dot{x}B - \frac{m}{\tau}\dot{y} + qE_y \\ 0 = -\frac{m}{\tau}\dot{z} + qE_z \end{cases}$$

Si l'on multiplie chacune des équations du système par nq , on fait apparaître $(nq\dot{x}, nq\dot{y}, nq\dot{z})$, soit (j_x, j_y, j_z) , d'où :

$$\begin{cases} 0 = qBj_y - \frac{m}{\tau}j_x + nq^2E_x \\ 0 = -qBj_x - \frac{m}{\tau}j_y + nq^2E_y \\ 0 = -\frac{m}{\tau}j_z + nq^2E_z \end{cases} \Rightarrow \begin{cases} j_x - \frac{q\tau B}{m}j_y = \frac{nq^2\tau}{m}E_x \\ \frac{q\tau B}{m}j_x + j_y = \frac{nq^2\tau}{m}E_y \\ j_z = \frac{nq^2\tau}{m}E_z \end{cases}$$

On en déduit la conductivité $\gamma = \frac{nq^2\tau}{m}$.

c) La résolution du système précédent pour le couple (j_x, j_y) donne les résultats demandés avec :

$$a = \frac{m^2}{m^2 + (q\tau B)^2} \quad \text{et} \quad b = \frac{q\tau Bm}{m^2 + (q\tau B)^2}$$

3. La conduction ne peut avoir lieu que suivant Ox , donc $j_y = 0$ et $j_z = 0$. Dans ce cas le système précédent devient :

$$\begin{cases} j_x = \gamma E_x \\ \frac{q\tau B}{m}j_x = \gamma E_y \\ 0 = \gamma E_z \end{cases}$$

On en déduit, qu'en plus du champ E_x qui assure la conduction (loi d'Ohm : $j_x = \gamma E_x$), il existe une composante E_y non nulle proportionnelle à B , c'est le champ de Hall. La mesure de la différence de potentiel créée par E_y entre les deux côtés du conducteur permet de mesurer B : c'est le principe d'une sonde à effet Hall.

La constante de Hall est :

$$R_H = \frac{q\tau}{m\gamma} = \frac{1}{nq}$$

B.3 1. Les deux forces qui s'exercent sur un électron sont : la force électrique et la force de frottement. On applique la relation fondamentale de la dynamique dans le référentiel galiléen :

$$m\vec{a} = q\vec{E} - \frac{m\vec{v}}{\tau}$$

2. On s'intéresse au régime forcé, donc après l'atténuation du régime transitoire. Mathématiquement, il s'agit de la solution particulière de l'équation différentielle. Ici le terme de forçage (c'est le champ $\vec{E}$) est sinusoïdal et suivant $\vec{u}_x$, donc on cherche une solution sinusoïdale pour $\vec{v}$ de même pulsation que le champ $\vec{E}$: on raisonne alors sur des grandeurs complexes. On rappelle que : $\frac{d\vec{v}}{dt} = i\omega \vec{v}$, donc l'équation du mouvement devient :

$$i\omega \vec{v} + \frac{\vec{v}}{\tau} = \frac{q}{m} \vec{E} \Rightarrow \vec{v} = \frac{\tau}{1 + i\omega\tau} \frac{q}{m} \vec{E}$$

Si l'on simplifie l'expression précédente par $\exp(i\omega t)$ il reste :

$$\vec{v}_0 e^{-i\varphi} = \frac{\tau}{1 + i\omega\tau} \frac{|q|}{m} E_0 \vec{u}_x$$

On en déduit pour la norme : $v_0 = \frac{qE_0}{m} \frac{\tau}{\sqrt{1 + (\omega\tau)^2}}$. En ce qui concerne φ , on a :

$$-\varphi = \text{Arg}\left(\frac{qE_0}{m} \frac{1}{1 + i\omega\tau}\right) = \pi - \varphi_D$$

puisque $\frac{qE_0}{m} < 0$ et si l'on note $\varphi_D = \text{Arg}(1 + i\omega\tau)$. Alors $\tan \varphi = \tan(-\pi + \varphi_D)$; soit $\tan \varphi = -\tan \varphi_D = -\omega\tau$.

3. On a donc :

$$\vec{j} = \vec{v} = \frac{\tau}{1 + i\omega\tau} \frac{nq^2}{m} \vec{E} \text{ soit } \underline{\gamma} = \frac{nq^2\tau}{m} \frac{1}{1 + i\omega\tau}$$

En continu, on a $\omega = 0$ soit $\gamma_0 = \gamma(0) = \frac{nq^2\tau}{m}$, d'où :

$$\underline{\gamma} = \frac{\gamma_0}{1 + i\omega\tau}$$

4. La fonction précédente correspond à une fonction passe-bas du premier ordre, de pulsation de coupure $\omega_c = \tau^{-1}$. La loi d'Ohm $\vec{j} = \gamma_0 \vec{E}$ sera vérifiée aux pulsations basses telles $\omega \ll \omega_c$. En revanche aux pulsations hautes, $\gamma \rightarrow 0$, et donc la loi n'est plus vérifiée : en fait la champ $\vec{E}$ varie trop vite, et les charges ont trop d'inertie pour se mettre en mouvement et "suivre" le champ.

B.4 1. L'électron se déplace à l'intérieur de la sphère donc il ressent le champ électrique créé par cette boule chargée pour une distance au centre inférieure au rayon a . On utilise le résultat de la question (3) de l'exercice (3). Pour une boule de rayon a , de densité volumique de charge ρ , le champ à l'intérieur est radial et vaut :

$$\vec{E}(r) = \frac{r\rho}{3\epsilon_0} \vec{u}_r$$

Dans le cas présent la boule a une charge e et un rayon a soit $\rho = e/(\frac{4}{3}\pi a^3)$, d'où :

$$\vec{E} = \frac{e}{4\pi\epsilon_0 a^3} r \vec{u}_r = \frac{e}{4\pi\epsilon_0 a^3} \vec{r}$$

puisque $\vec{r} = r\vec{u}_r$.

La force ressentie par l'électron est donc :

$$\vec{F} = -e\vec{E} = \frac{-e^2}{4\pi\epsilon_0 a^3} \vec{r} = -k \vec{r}$$

2. a) Il s'agit d'une force centrale donc le mouvement est plan. On rappelle que ceci se démontre en calculant la dérivée du moment cinétique par rapport à O $\vec{\sigma}_O$ avec le théorème du moment cinétique :

$$\frac{d\vec{\sigma}_O}{dt} = \vec{r} \wedge \vec{F} = \vec{0}$$

Le moment cinétique est une constante donc $\vec{r}$ et la vitesse $\vec{v}$ sont dans un plan perpendiculaire à $\vec{\sigma}_O$. Le mouvement a lieu le plan perpendiculaire à $\vec{\sigma}_O$ contenant la position initiale de M .

- b) On applique la relation fondamentale de la dynamique à l'électron et l'on projette sur les trois axes :

$$m_e \frac{d^2 \vec{r}}{dt^2} = -k \vec{r} \Rightarrow \begin{cases} m_e \ddot{x} + kx = 0 \\ m_e \ddot{y} + ky = 0 \\ m_e \ddot{z} + kz = 0 \end{cases} \Rightarrow \begin{cases} \ddot{x} + \omega_0^2 x = 0 \\ \ddot{y} + \omega_0^2 y = 0 \\ \ddot{z} + \omega_0^2 z = 0 \end{cases} \quad \text{avec} \quad \omega_0 = \sqrt{\frac{k}{m_e}}.$$

- c) La solution du système précédent est :

$$\begin{cases} x = A_1 \cos \omega_0 t + A_2 \sin \omega_0 t \\ y = B_1 \cos \omega_0 t + B_2 \sin \omega_0 t \\ z = C_1 \cos \omega_0 t + C_2 \sin \omega_0 t \end{cases} \quad \text{soit} \quad \begin{cases} \dot{x} = \omega_0 (-A_1 \sin \omega_0 t + A_2 \cos \omega_0 t) \\ \dot{y} = \omega_0 (-B_1 \sin \omega_0 t + B_2 \cos \omega_0 t) \\ \dot{z} = \omega_0 (-C_1 \sin \omega_0 t + C_2 \cos \omega_0 t) \end{cases}$$

Les conditions initiales donnent :

$$\begin{cases} x(0) = A_1 = r_0 \\ y(0) = B_1 = 0 \\ z(0) = C_1 = 0 \end{cases} \quad \text{et} \quad \begin{cases} \dot{x}(0) = \omega_0 A_2 = 0 \\ \dot{y}(0) = \omega_0 B_2 = 0 \\ \dot{z}(0) = \omega_0 C_2 = v_0 \end{cases}$$

On en déduit que $y(t) = 0$, le mouvement a lieu dans le plan xOz et $x(t) = r_0 \cos \omega_0 t$, et

$$z(t) = \frac{v_0}{\omega_0} \sin \omega_0 t. \text{ La trajectoire est une ellipse d'équation : } \left(\frac{x}{r_0} \right)^2 + \left(\frac{\omega_0 z}{v_0} \right)^2 = 1$$

- d) La période est $T_0 = \frac{2\pi}{\omega_0} = 2\pi \sqrt{\frac{m_e}{k}}$.

- B.5** 1. a) On fait le bilan des forces s'exerçant sur la gouttelette dans le référentiel galiléen : le poids, la poussée d'Archimède (opposée du poids du volume d'air déplacé) et le frottement. La relation fondamentale de la dynamique (R.F.D.) appliquée à une goutte donne :

$$m \vec{a} = \frac{4}{3} \pi R^3 \rho_h \vec{g} - \frac{4}{3} \pi R^3 \rho_a \vec{g} - k \vec{v}$$

La vitesse limite est obtenue lorsque l'accélération devient nulle :

$$\vec{v}_0 = \frac{4}{3} \pi R^3 \frac{1}{k} (\rho_h - \rho_a) \vec{g} \simeq \frac{4}{3} \pi R^3 \frac{1}{k} \rho_h \vec{g} \simeq \frac{4}{3} \pi R^2 \frac{1}{\alpha} \rho_h \vec{g}$$

- b) On peut réécrire la R.F.D avec $\vec{v}_0$:

$$m \frac{d\vec{v}}{dt} = k(\vec{v}_0 - \vec{v}) \quad \text{ou} \quad \frac{d\vec{v}}{dt} + \frac{\vec{v}}{\tau} = \frac{\vec{v}_0}{\tau}$$

où l'on a défini le temps caractéristique $\tau = m/k$.

La solution de cette équation différentielle est :

$$\vec{v}(t) = \vec{A} \exp\left(-\frac{t}{\tau}\right) + \vec{v}_0$$

À $t = 0$ la vitesse est nulle soit $\vec{0} = \vec{A} + \vec{v}_0$; on en déduit l'expression finale de $\vec{v}$:

$$\vec{v}(t) = \left(1 - \exp\left(-\frac{t}{\tau}\right)\right) \vec{v}_0$$

- c) Application numérique : avec l'expression de $\vec{v}_0$ on trouve $R = 1,13 \cdot 10^{-6}$ m.

2. a) Le champ électrique est suivant les potentiels décroissants $\vec{E} = -E \vec{u}_z$, donc dans ce cas $V_1 > V_2$ et $U > 0$. Puisque $\vec{E} = -\vec{\text{grad}} V$, on a en projection sur zz' :

$$\frac{dV}{dz} = E \Rightarrow \int_{V_2}^{V_1} dV = \int_{z_2}^{z_1} E dz \Rightarrow U = V_1 - V_2 = E(z_1 - z_2) = Ed$$

z_1 et z_2 étant les positions des plaques.

- b) La force électrique s'exerçant sur la goutte est $\vec{F} = -q\vec{E}$. À l'équilibre, la force de frottement est nulle, donc la force électrique équilibre le poids et la poussée d'Archimède :

$$\vec{0} = -q\vec{E} + \frac{4}{3}\pi R^3(\rho_h - \rho_a)\vec{g} \Rightarrow 0 = q\frac{U}{d} - \frac{4}{3}\pi R^3(\rho_h - \rho_a)g$$

d'où :

$$q = \frac{d}{U} \frac{4}{3}\pi R^3(\rho_h - \rho_a)g \Rightarrow q = 4,8 \cdot 10^{-19} \text{ C}$$

B.6 1. a) La seule force qui s'applique est la force magnétique $\vec{F}_B = q(\vec{v} \wedge \vec{B})$. Pour un déplacement élémentaire $d\vec{\ell}$, le travail de cette force est $\delta W = \vec{F} \cdot d\vec{\ell} = \vec{0}$ puisque $\vec{v}$ et $d\vec{\ell}$ sont colinéaires. L'énergie cinétique est donc constante.

- b) La relation fondamentale de la dynamique donne :

$$m \frac{d\vec{v}}{dt} = q(\vec{v} \wedge \vec{B})$$

La projection de la force $\vec{F}_B$ sur Oz est nulle puisqu'elle est perpendiculaire à $\vec{B}$ qui est selon Oz . La projection de la relation fondamentale sur Oz donne donc :

$$m \frac{dv_L}{dt} = 0 \Rightarrow \vec{v}_L = c\vec{e}$$

L'énergie cinétique est constante donc la norme de la vitesse $\|\vec{v}\| = \sqrt{v_L^2 + v_{\perp}^2}$ est aussi constante. Or v_L est elle-même constante donc $v_{\perp}$ l'est aussi.

- c) On projette la relation fondamentale dans le plan xOy :

$$\vec{v} \wedge \vec{B} \begin{cases} Bv_y \\ -Bv_x \end{cases} \text{ soit } \begin{cases} mv_x = qBv_y \\ -mv_y = qBv_x \end{cases} \Rightarrow \begin{cases} \dot{v}_x = \omega v_y \\ \dot{v}_y = -\omega v_x \end{cases}$$

En reportant v_y de la première équation dans la deuxième, on obtient : $\ddot{v}_x + \omega^2 v_x = 0$, dont la solution est $v_x(t) = A_1 \cos \omega t + A_2 \sin \omega t$. À $t = 0$, $v_x(0) = v_{\perp 0} = A_1$.

La première équation du système nous donne donc :

$$\dot{v}_x = -\omega v_{\perp 0} \sin \omega t + \omega A_2 \cos \omega t = \omega v_y$$

soit $v_y = -v_{\perp 0} \sin \omega t + A_2 \cos \omega t$, or à $t = 0$, $v_y = A_2 = 0$. On a donc :

$$\begin{cases} v_x = v_{\perp 0} \cos \omega t \\ v_y = -v_{\perp 0} \sin \omega t \end{cases}$$

- d) On intègre le système précédent :

$$\begin{cases} v_x = v_{\perp 0} \cos \omega t \\ v_y = -v_{\perp 0} \sin \omega t \end{cases} \Rightarrow \begin{cases} x = \frac{v_{\perp 0}}{\omega} \sin \omega t + c t e_1 \\ y = \frac{v_{\perp 0}}{\omega} \cos \omega t + c t e_2 \end{cases}$$

Sachant qu'à $t = 0$ on a $x = 0$ et $y = 0$, on en déduit :

$$\begin{cases} x = \frac{v_{\perp 0}}{\omega} \sin \omega t \\ y = \frac{v_{\perp 0}}{\omega} (\cos \omega t - 1) \end{cases} \quad \text{et} \quad x^2 + \left(y - \frac{v_{\perp 0}}{\omega}\right)^2 = \left(\frac{v_{\perp 0}}{\omega}\right)^2$$

Si l'on note $a = \frac{v_{\perp 0}}{\omega}$, l'équation précédente est celle d'un cercle de rayon a et de centre $C : (x_c = 0, y_c = a)$.

La période de révolution est $T_1 = 2\pi/\omega$

- e) La particule de charge q décrit un tour en un temps T_1 , donc $i = q/T_1$. D'après les équations paramétriques, la particule tourne dans le sens horaire, donc la normale au cercle (règle de la main droite, ou du « tire bouchon ») est selon $-\vec{u}_z$, d'où la définition de $\vec{\mu}$. On en déduit :

$$\vec{\mu} = -\frac{\pi a^2 q}{T_1} \vec{u}_z. \text{ Ainsi :}$$

$$\mu = \frac{\pi a^2 q \omega}{2\pi} \text{ soit avec } a = \frac{v_{\perp 0}}{\omega} : \mu = \frac{v_{\perp 0}^2 q}{2\omega} = \frac{v_{\perp 0}^2 q}{2} \frac{m}{qB} = \frac{E_{c\perp}}{B}$$

L'autre expression demandée est obtenue avec $\varphi = BS = B\pi a^2$:

$$\mu = \frac{\varphi q^2}{2\pi m}$$

- f) La vitesse $\vec{v}_L$ selon $\vec{u}_z$ est constante et la projection dans le plan xOy de la trajectoire est un cercle. Le mouvement de la particule résulte de la composition d'un mouvement circulaire perpendiculairement à Oz et d'une translation selon Oz , la trajectoire est donc une hélice de pas régulier (hauteur parcourue en un tour), d'axe parallèle à Oz et passant par C . La particule s'enroule donc sur un cylindre d'axe Oz parallèle à $\vec{B}$ uniforme, ce cylindre constitue donc un tube de champ (par définition, tube dont les génératrices sont des lignes de champs).

En un temps T_1 , le centre C se déplace de $b = v_L T_1$ selon $\vec{u}_z$, puisque v_L est constante, soit $b = \frac{2\pi v_L}{\omega}$.

2. a) D'après l'énoncé, on s'intéresse à un champ qui a la symétrie de révolution, donc $\vec{B}$ ne dépend pas de θ . On applique le théorème d'Ampère sur un contour circulaire de centre O' (figure S52.2), de rayon r et d'axe Oz :

$$\oint \vec{B} \cdot d\vec{\ell} = \oint \vec{B} \cdot \vec{u}_\theta d\ell = \oint B_\theta(r, z) r d\theta$$

Comme B_θ ne dépend pas de θ on peut le sortir de l'intégrale qui se réduit alors à $2\pi r B_\theta$.

Si on note I_{enl} le courant enlacé par le contour C , le théorème d'Ampère donne : $2\pi r B_\theta = \mu_0 I_{\text{enl}}$. D'après l'énoncé, il n'y a pas de courant dans la zone considérée donc $I_{\text{enl}} = 0$ et $B_\theta = 0$.

- b) Le champ $\vec{B}$ est à flux conservatif, ce qui signifie que le champ de $\vec{B}$ à travers une surface fermée est nul. Comment choisir la surface : on cherche la relation entre les composantes B_r et B_z , on choisit donc comme surface fermée Σ , un petit cylindre de hauteur dz et de rayon r . Il ne faut pas oublier que par convention, les vecteurs surfaces sont orientés vers l'extérieur pour une surface fermée.

Le cylindre est de petit rayon donc on peut considérer que sur les surfaces supérieures et inférieures, le champ est sensiblement le même que sur l'axe $B_z(0, z)$ et $B_z(0, z + dz)$. De même, sur la surface latérale, on supposera que le champ a quasiment la même valeur qu'à la cote z . Le flux $\iint_{\Sigma} \vec{B} \cdot d\vec{S}$

se décompose en trois termes :

A. Flux à travers la surface supérieure :

$$\iint_{S(z+dz)} \vec{B} \cdot d\vec{S}_{z+dz} = \iint_{S(z+dz)} \vec{B} \cdot \vec{u}_z dS_{z+dz} = \iint_{S(z+dz)} B_z dS_{z+dz} = \pi r^2 B_z(0, z + dz)$$

B. Flux à travers la surface inférieure :

$$\iint_{S(z)} \vec{B} \cdot d\vec{S}_z = \iint_{S(z)} \vec{B} \cdot (-\vec{u}_z) dS_z = - \iint_{S(z)} B_z dS_z = -\pi r^2 B_z(0, z)$$

C. Flux à travers la surface latérale :

$$\iint_{S(L)} \vec{B} \cdot d\vec{S}_L = \iint_{S(L)} \vec{B} \cdot \vec{u}_r dS_L = \iint_{S(L)} B_r dS_L = 2\pi r dz B_r(r, z)$$

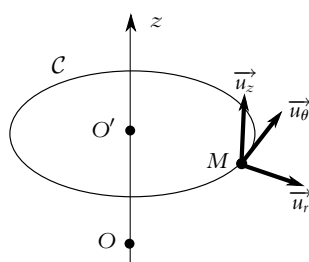


Figure S52.2

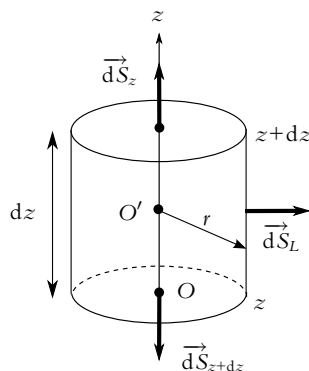


Figure S52.3

La somme de ces trois termes est nulle soit :

$$\pi r^2 (B_z(0, z + dz) - B_z(0, z)) + 2\pi r B_r(r, z) dz = 0$$

Or par un développement au premier ordre : $B_z(0, z + dz) - B_z(0, z) = \frac{dB_z}{dz} dz$, d’où :

$$B_r = -\frac{r}{2} \frac{dB_z}{dz}$$

c) On calcule la force magnétique en coordonnées cylindriques avec $B_\theta = 0$:

$$\vec{F} = q \vec{v} \wedge \vec{B} \begin{cases} q v_\theta B_z \\ q(v_z B_r - v_r B_z) \\ -q v_\theta B_r \end{cases}$$

Donc $F_z = -q v_\theta B_r$, or on a montré dans la première partie que la particule tournait dans le sens rétrograde donc $v_\theta < 0$. D’autre part, d’après les hypothèses, la particule décrit une trajectoire circulaire dans les plans perpendiculaires à Oz donc $v_r = 0$ et ainsi $|v_\theta| = v_\perp$ soit $v_\theta = -v_\perp$. D’après les hypothèses encore et les résultats de la première partie, le rayon r est égal à $\frac{mv_\perp}{qB_z}$. Enfin en utilisant, le résultat de la question précédente :

$$F_z = -q v_\theta B_r = q v_\perp B_r = q v_\perp \left(-\frac{r}{2} \frac{dB_z}{dz} \right) = -\frac{mv_\perp^2}{2B_z} \frac{dB_z}{dz}$$

d) On calcule $\frac{d\mu}{dt}$ à partir de l’expression $\mu = \frac{E_{c\perp}}{B_z}$ où l’on a remplacé B par B_z :

$$\frac{d\mu}{dt} = \frac{1}{B_z} \frac{dE_{c\perp}}{dt} - \frac{E_{c\perp}}{B_z^2} \frac{dB_z}{dt} = \frac{1}{B_z} \frac{dE_{c\perp}}{dt} - \frac{E_{c\perp}}{B_z^2} \frac{dB_z}{dz} \frac{dz}{dt}$$

or $\frac{dz}{dt} = v_L$ et $\frac{E_{c\perp}}{B_z^2} \frac{dB_z}{dz} = -\frac{1}{B_z} F_z$. De plus $F_z v_L$ est la puissance de F_z , donc :

$$\frac{d\mu}{dt} = \frac{1}{B_z} \left(\frac{dE_{c\perp}}{dt} + F_z v_L \right) = \frac{1}{B_z} \left(\frac{dE_{c\perp}}{dt} + \frac{dE_c}{dt} \right) = \frac{1}{B_z} \frac{dE_c}{dt} = 0$$

puisque la force magnétique ne travaille pas $E_c = te$.

On rappelle la deuxième expression de μ faisant intervenir le flux de $\vec{B}$: $\mu = \frac{eq^2}{2\pi m}$ à travers la trajectoire de la particule. Ainsi, si μ est constant, le flux de $\vec{B}$ l’est aussi, or le fait que le flux soit le même est l’une des propriétés de toute section d’un tube de champ (cela se redémontre en utilisant le fait que le flux de $\vec{B}$ à travers une surface fermée est nul, et que le flux à travers la surface formée par les génératrices du tube est nul aussi puisque ces génératrices sont des lignes de champ). On peut donc considérer que la particule s’enroule sur un tube de champ.

e) On écrit l’énergie cinétique : $E_c = \frac{1}{2}m(v_L^2 + v_\perp^2)$ ou $E_c = \frac{1}{2}mv_L^2 + E_{c\perp}$; soit avec l’expression de μ :

$$E_c = \frac{1}{2}mv_L^2 + \mu B_z$$

f) On a déjà montré que E_c est constante ; μ l’est aussi et $mv_L^2 \geq 0$, on en déduit que le mouvement ne peut avoir lieu que dans les zones où $E_c \geq \mu B_z$, soit $B_z \leq B_{\max}(= E_c/\mu)$

g) De O à M_1 ainsi que de O à M'_1 , les lignes de champ se resserrent, donc l’intensité du champ augmente. Cela se démontre en utilisant une propriété des tubes de champ précitée, que le flux à travers toute section d’un tube est le même : si la section a une surface plus petite cela correspond donc à un champ plus intense.

Le plan xOy est un plan d’antysymétrie pour le champ $\vec{B}$, c’est à dire que pour deux points M et M' symétriques par rapport à xOy , on a $\vec{B}(M') = -\text{sym}(\vec{B}(M))$, où $\text{sym}(\vec{B}(M))$ est le vecteur symétrique de $\vec{B}(M)$ par le plan xOy . Le intensités des champs sont donc les mêmes en M_1 et M'_1 .

h) On suppose donc que le champ $B_{\max}$ est atteint en M_O et m'_O , respectivement avant M_1 et M'_1 , donc d’après la réponse (f), la particule reste confinée dans la zone où $B_z \leq B_{\max}$, c’est-à-dire entre M_0 et M'_0 .

Le centre guide C , se déplace à la vitesse $\frac{dz}{dt} = v_L$, soit

$$v_L^2 = \frac{2}{m}(E_c - \mu B_z) = \frac{2\mu}{m}(B_{\max} - B_z) \Rightarrow dt = \pm \sqrt{\frac{m}{2\mu}} \frac{dz}{\sqrt{B_{\max} - B_z}}$$

La période correspond à un trajet de O à M_O , puis de M_O à M'_O et enfin de M'_O à O , soit quatre trajet entre $z = 0$ et $z = z_0$, soit :

$$T_2 = 4 \pm \sqrt{\frac{m}{2\mu}} \int_0^{z_0} \frac{dz}{\sqrt{B_{\max} - B_z}}$$

B.7 1. $\vec{V} = -\frac{1}{ne} \vec{J}$. Le champ magnétique dévie dans un premier temps les électrons vers la face A' . Il apparaît donc un excès d'électrons sur la face A' et un déficit sur la face A , ce qui se traduit par l'existence d'un champ électrique $\vec{E}_H$, dirigé de la face A vers la face A' , dont l'action en régime permanent compense exactement celle du champ magnétique. On en déduit : $\vec{E}_H + \vec{V} \wedge \vec{B} = 0$ d'où $\vec{E}_H = \frac{1}{ne} \vec{J} \wedge \vec{B} = -\frac{I}{nebh} B \vec{u}_y$.

2. $U_H = \int_1^{1'} \vec{E} \cdot (-dy \vec{u}_y) = \frac{IB}{neh} = \frac{C_H}{h} IB$ avec $C_H = \frac{1}{ne}$.

A.N. : $U_H = 125$ mV et $n = 1,7 \cdot 10^{22}$ électrons par mètre cube.

σ représente la conductivité du matériau. En présence du champ magnétique, le champ électrique est $\vec{E} = \vec{E}_H + \vec{E}'$ donc $\vec{J} = \sigma \vec{E}' = \sigma (\vec{E} - C_H \vec{J} \wedge \vec{B})$.

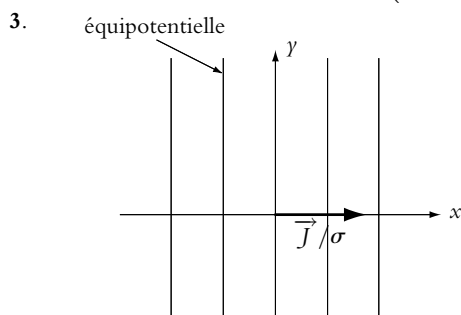


Figure S52.4 En l'absence de champ magnétique

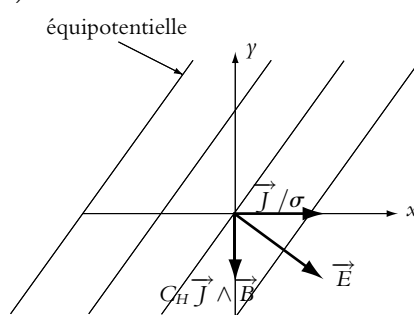
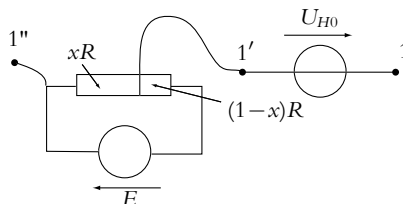


Figure S52.5 En présence de champ magnétique

4. On voit sur la figure S52.5 que $\tan \theta = \frac{C_H JB}{J/\sigma} = \frac{\sigma}{ne} B$ donc θ ne dépend que du semi-conducteur et de B . Pour le semi-conducteur étudié, θ est compris entre $-\pi/2$ et 0 .
5. Si $U_{H0} \neq 0$, il faut brancher en série sur le montage une source de tension de f.e.m. réglable grâce à un potentiomètre pour compenser cette tension résiduelle. Les bornes du montage sont maintenant 1 et 1" au lieu de 1 et 1'.



Index



A

aberration
 chromatique 268
 de coma 270
accélération
 absolue 646
 complémentaire ou de Coriolis 650
 d'entraînement 646, 650
 relative 646
admittance complexe 338
agitation thermique 788
alimentation stabilisée 481
amplificateur
 inverseur 536
 opérationnel 508
analyse
 de Fourier 447
 spectrale 453
aplanétisme rigoureux 221
apogée 743
approximation
 de l'optique géométrique 201
autocollimation 289
axe optique 203

B

balayage 487
 en fréquence 481
bande
 passante 402, 495, 616
 à -3 dB 360
 rejetée 434
base de temps 483
biréfringent 201
bobine de Helmholtz 1118
bonnette 292

C

calorimètre 851
calorimétrie 851
capacité thermique
 à pression constante 848
 massique du liquide saturant seul 927

caractéristique de transfert 510
centre
 d'inertie 678
 de gravité 678
 du miroir 228
 optique 248
chaleur latente de changement d'état 925

champ
 à flux conservatif 1100
 angulaire 279
 de gradient 1006
 de pesanteur 763
 électrostatique 979
 en profondeur 279
 gravitationnel 760
 magnétique 1086
 scalaire 978
 stationnaire 979
 uniforme 979
 vectoriel 978

changement
 de phase 914
 de référentiel 644

charge
 électrique 966
 ponctuelle 970
circulation 1003
cloison diatherme 881
coefficient
 de compressibilité isotherme 799
 de dilatation isobare 799
 de surtension 363
 thermoélastique 799

colatitude 600
collimateur 293
comparateur 539
 à hystérésis 540
 simple 539

composante
 axiale 598
 continue 453
 orthoradiale 598
 radiale 598, 602

condition(s)
 d'émergence 315
 de Gauss 227
conductivité 1163
cône de réfraction 209
conformateur à diodes 548
conjugué 203
constante
 de Boltzmann 797
 de Hall 1181
 des aires 710
 des gaz parfaits 796
contour d'Ampère 1111
coordonnées
 cartésiennes 88, 595
 cylindriques 596
 polaires 89
 sphériques 599
courant
 de conduction 1079
 de convection 1079
 de diffusion 1079
 de polarisation 521
 électrique 1078

courbe
 d'ébullition 919
 de rosée 919
 isotitre 921
covolume 798
création d'entropie 880
cycle
 de Carnot 858
 de Hirn 956
 ditherme 858
 moteur 845

D

décade 407
demi-grand axe 740
démodulation synchrone 581
densité
 de courant 1081
 linéique
 de charges 969

de courant 1083
 surfacique
 de charges 969
 de courant 1083
 volumique
 de charges 968
 de courant 1082
 déphasage 457
 dérivée d'un vecteur unitaire
 tournant 604
 détecteur de crête 576
 détente 861
 de Joule - Gay Lussac 861
 de Joule - Kelvin 864
 de Joule - Thomson 864
 deuxième principe 879
 déviation vers l'Est 775
 diagramme
 de Bode 401
 de Clapeyron 845
 de Raveau 942
 de Watt 845
 diffraction 198
 diode
 à jonction 563
 bloquée ou bloquante 564
 passante 564
 polarisée
 en direct 564
 en inverse 564
 dioptre 203
 plan 224
 dioptrie 249
 dipôle
 électrostatique 1053
 magnétique 1138
 polarisé 564
 distance
 focale 740
 image 229
 objet 229
 distorsion 266
 distribution
 de charges 966
 dipolaire 1069
 quadripolaire 1069
 unipolaire 1068

E

échelle
 Celsius 794
 centésimale linéaire 791
 légale 794
 Echelle Internationale de
 Température (EIT)
 794

effet Hall 1179
 efficacité 945
 élément cinétique 680
 émission stimulée 284
 énergie
 électrostatique 1012
 potentielle 1010, 1011
 d'interaction 1012
 d'interaction entre deux
 charges 1012
 d'un dipôle dans $\vec{B}$ 1144
 d'un dipôle dans $\vec{E}$ 1063
 effective 711
 enrichissement du spectre 577
 enthalpie 847
 massique
 de changement d'état
 925
 de fusion 853
 entrée
 inverseuse 508
 non inverseuse 508
 entropie 879
 statistique 909
 équation(s)
 barométrique 806
 canonique d'un filtre passe-
 bas du
 premier ordre 406
 d'état 796
 de Laplace 857
 fondamentale de la statique
 des
 fluides 805
 équilibre thermodynamique 788
 équipotentielle 1016
 espace
 image 204
 objet 204
 état
 fluide 917
 métastable 915
 macroscopique le plus pro-
 bable 906
 évolution
 ditherme 896
 monotherme 896
 excentricité 740

F

facteur
 d'échelle 421
 de Boltzmann 811
 de forme 421
 de puissance 383
 de qualité 619

faisceau
 convergent 205
 divergent 205
 lumineux 197
 parallèle 205
 filtrage 447
 filtre
 coupe-bande du second
 ordre 431
 du premier ordre 404
 idéal 458
 passe-bande 426, 459
 passe-bas 459
 du premier ordre 405
 du second ordre 418
 passe-haut 459
 du premier ordre 410
 du second ordre 422
 sélectif 429, 468
 fluide 786
 caloporteur 946
 homogène 805
 incompressible 805
 flux 998
 élémentaire 998
 sortant 999
 fonction
 de transfert 397
 complexe 397
 d'état 836
 fondamental 453
 force
 centrale 721
 centrifuge 665
 conservative 148
 contre électromotrice 567
 d'inertie
 complémentaire ou de
 Coriolis 664
 d'entraînement 664
 de Laplace 1095
 de Lorentz 1086, 1152
 extérieure 684
 intérieure 684
 formule(s)
 de Binet 732
 de conjugaison des lentilles
 250
 de conjugaison du miroir
 232
 foyer
 image 221
 secondaire 222
 objet 221
 secondaire 222
 fréquence 200

Index

G

gain 399, 457
 en décibels 399
 en puissance 399
 en tension 400
Gauss 1087
gaz 786
 hypercritique 917
générateur(s)
 basses fréquences ou GBF 478
goniomètre 293
gradient 1004
grandissement 223
 angulaire 240
grossissement 240

H

harmonique d'ordre n 453
hypothèse microcanonique 906

I

identité thermodynamique 884
image 203
 réelle 205
 virtuelle 205
impédance
 d'entrée 396, 494, 522
 de sortie 396, 522
 mécanique 619
indice de réfraction 199
inégalité de Clausius 896
intensité
 électrique 1080
interaction
 attractive 727
 gravitationnelle 759
 newtonienne 726
 répulsive 727
interférence 198
invariance 984
invariant de Runge-Lenz ou de Laplace 731
isotherme d'Andrews 919

L

lampe à incandescence 282
Laser 284
latitude 761
lentille
 à bords épais 247
 à bords minces 247
 convergente 248
 divergente 248
 sphérique 247
 sphérique mince 248
ligne

 de champ 1016
 de courant 1082
limite de Roche 773
liquide 786, 913
 saturant 919
loi(s)
 d'Ohm
 intégrale 1164
 locale 1163
 de Biot et Savart 1092
 de Cauchy 314
 de composition des accélérations 650
 de composition des vitesses 647
 de Coulomb 976
 de Gladstone 211
 de Kepler 736
longitude 761
 ou azimut 600
longueur d'onde 200
loupe 285
lunette
 à frontale fixe 287
 afocale 287
 astronomique 263
 de Galilée 262
 de visée
 à l'infini 287
 autocollimatrice 289

M

machine
 ditherme 940
 frigorifique 945
 monotherme 939
 thermique 938
macroétat 906
macroscopique 786
magnéton
 de Bohr 1139
 nucléaire 1139
marée
 basse 771
 de « vives eaux » 772
 de « mortes eaux » 772
 d'équinoxe 772
 haute 771
Masse 490
masse réduite 705
méridien 600, 761
mésoscopique 786
méthode
 de Bessel 310
 de Lissajous 500
 de Silberman 309

 des points conjugués 311
microétat 906
microscope 260
microscopique 786
milieu
 homogène 200, 201, 913
 isotrope 200, 201
minimum de déviation 316
mirage 212
miroir 203
 concave 228
 convexe 228
 plan 222
 sphérique 228
mise au point 289
mode
 AC 486
 DC 486
 XY 487
modèle
 de Drude 1162
 idéal de l'amplificateur opérationnel 511
modulation d'amplitude 579
moment
 cinétique 679
 barycentrique 682
 par rapport à un axe 632
 par rapport à un point 631
 de la force $\vec{f}$ par rapport à un point O 633
dipolaire 1053
 induit 1070
 permanent 1070
 quadrupolaire 1070
 résultant 685
moment m d'une force par rapport à un axe 634
monochromatique 200
monovariant 916
montage
 dérivateur 538
 sommateur 537
mouvement relatif 703
multimètre 493
multiplieur 578
multivibrateur astable 544

N

nombre d'Avogadro 787

O

objectif 286
 objet
 réel 205
 virtuel 205
 octave 408
 oculaire 285
 œil 278
 astigmat 281
 emmétrope 280
 hypermétrope 281
 myope 280
 presbyte 281
 ondulation 380
 opalescence critique 922
 ordre du filtre 460
 orientation d'une surface 996
 oscillateur
 amorti 176
 de relaxation 478
 harmonique 171
 mécanique 610
 oscilloscope 482
 analogique 482
 numérique 484

P

palier de changement d'état 919
 parallaxe 289
 parallèle 600, 761
 particule
 de fluide 803
 fictive 706
 périgée 743
 perméabilité du vide 1093
 permittivité diélectrique du vide 977
 phase 913
 pinceau lumineux 197
 plan
 d'antisymétrie 987
 d'incidence 207
 de symétrie 987
 focal
 image 221
 objet 221
 méri dien 600
 polaire 597
 point
 coïncidant 645
 critique 917
 d'incidence 207
 triple 916
 polarisabilité 1071
 pompe à chaleur 947
 pont de Gaetz 573

porteuse 579
 potentiel
 électrostatique 1006
 newtonien 726
 pouvoir séparateur 279
 premier principe 835
 premier théorème de Koenig 682
 première vitesse cosmique 743
 pression 789
 de vapeur saturante 922
 moléculaire 798
 thermodynamique 883
 principe
 de Carnot 940
 de Curie 984
 de relativité galiléenne 662
 de superposition 977, 1087
 prisme 313
 pseudo-vecteur 989, 1089
 puissance
 active 383
 des forces intérieures 691
 moyenne 379
 pulsation
 cyclotron 1166
 de coupure 403
 Punctum Proximum 279
 Punctum Remotum 279

Q

quadripôle 395, 578

R

rayon
 algébrique 228
 émergent 204
 incident 204
 lumineux 197
 paraxial 227
 vert 212
 redressement 569
 double-alternance 569
 monoalternance 569
 référentiel
 absolu 645
 barycentrique 680
 de Copernic 661
 de Kepler 661
 géocentrique 661
 relatif 645
 terrestre 660
 réflexion 206
 totale 209
 réfraction 207
 limite 209

régime

 forcé 612
 libre 612
 linéaire 510
 saturé 510
 transitoire 612
 relation(s)
 de Clausius 896
 de Descartes 252
 de Lagrange-Helmholtz 240
 de Mayer 848
 de Newton 231, 251
 rendement 943
 repères thermométriques 791
 résistance dynamique ou diffé-
 rentielle 566
 résonance
 aiguë 361
 en amplitude 615
 en intensité 357
 en vitesse 618
 floue 361
 résultante
 cinétique 678
 des forces 685
 retard au changement d'état 915
 réticule 288
 retour inverse de la lumière 202
 RMS 493

S

satellite géostationnaire 747
 saturation en courant 515
 second théorème de Koenig 683
 seconde loi de Joule 848
 seconde vitesse cosmique 745
 série de Fourier 449
 signal modulant 579
 solide 786, 913
 sommet du miroir 228
 source
 de chaleur 888
 de déclenchement 488
 spectre 282
 de Fourier 454
 spectromètre à prisme 318
 stigmatisme
 approché 223
 rigoureux 221
 surface
 de Gauss 1027
 fermée 998
 surfusion 915
 symétrie 984
 synchronisation 487
 système

Index

- afocal 223
 - catadioptrique 203
 - de forces 684
 - dioptrique 203
 - divariant 916
 - fermé 836
 - optique 203
 - ouvert 836
 - thermodynamique 786
- T**
- taux de modulation 579
 - température 790
 - absolue 793
 - empirique 791
 - thermodynamique 881
 - tension
 - d'offset 479
 - de décalage ou offset 518
 - de Hall 1180
 - de saturation 510, 514
 - de seuil 567
 - terme de marée 766
 - Terre 490
 - tesla 1087
 - théorème
 - d'Ampère 1102
 - de Carnot 944
 - du moment cinétique 634
 - thermostat 888
 - titre
 - massique 919
 - molaire 919
 - transadmittance complexe 398
 - transconductance 397
 - transfert
 - statique 396
 - thermique 837
 - transformation
 - adiabatique 839
 - de Galilée 660
 - isenthalpique 866
 - isobare 839, 846
 - isochore 839
 - isotherme 839
 - monobare 846
 - quasi-statique 840
 - réversible 841, 878
 - transimpédance complexe 398
 - transition de phase 914
 - transrésistance 397
 - travail 837
 - maximal récupérable 939
 - Trigger ou « Déclencheur de Schmidt » 543
 - troisième loi de Kepler 742
 - troisième principe 911
 - troisième vitesse cosmique 745
 - TTL 478
 - tube de courant 1082
- V**
- valeur
 - efficace 381, 493
 - en eau du calorimètre 851
 - moyenne 379, 453, 493
 - Van der Waals 797
 - vapeur 913
 - saturante 919
 - sèche 919
 - variable
 - d'état 795
 - extensive 795
 - intensive 795
 - variété allotropique 914
 - vecteur
 - axial 989, 1089
 - excentricité 729
 - polaire 989
 - rotation 607
 - vergence 249
 - viseur 287
 - vitesse
 - absolue 646
 - aréolaire 710
 - d'entraînement 646, 648
 - de balayage ou slew rate 518
 - de libération 745
 - relative 646
- W**
- wobulation 481
- Z**
- zone d'avalanche 565

sous la direction de Marie-Noëlle Sanz
Anne-Emmanuelle Badel • François Clausset

PHYSIQUE

TOUT-EN-UN • 1^{re} ANNÉE

Cours et exercices corrigés

3^e édition

Cet ouvrage propose aux étudiants de première année MPSI, PCSI et PTSI un cours complet et 380 exercices et problèmes intégralement corrigés.

Cette nouvelle édition, en deux couleurs, est entièrement revue et corrigée. Les exercices et problèmes sont plus nombreux et en grande partie renouvelés

Un cours complet et conforme au programme

- Toutes les notions du programme sont abordées avec clarté et concision.
- De nombreuses illustrations facilitent la bonne assimilation du cours.

380 exercices et problèmes de synthèse

Dans chaque chapitre :

- des « applications directes du cours » pour tester ses connaissances,
- des « exercices et problèmes » pour s'entraîner et préparer les concours.

Toutes les solutions détaillées

À la fin de l'ouvrage sont regroupées les solutions détaillées des exercices et des problèmes.

Des exercices supplémentaires et leurs corrigés sont disponibles sur www.dunod.com

MARIE-NOËLLE SANZ
est professeur en PC* au
lycée Saint-Louis à Paris.

ANNE-EMMANUELLE
BADEL

est professeur en
BCPST au lycée
du Parc à Lyon.

FRANÇOIS CLAUSSET
est professeur en MP
au lycée Jean Perrin
à Lyon.

MATHÉMATIQUES

PHYSIQUE

CHIMIE

SCIENCES DE LA VIE

SCIENCES DE LA TERRE

